

全国最低交易手续费免费办理国内正规商品股指原油期货开户 零佣金 不限资金量和交易量直接申请交易所手续费
任何地方费用比我们低 10倍赔偿差价 客服QQ/微信:958218088 期货开户和书籍下载请进:https://www.7help.net



CFA
INSTITUTE

CFA协会投资系列

CFA协会
推荐教材

QUANTITATIVE INVESTMENT
ANALYSIS (Second Edition)

定量投资分析

(美) 理查德 A. 德弗斯科 丹尼斯 W. 麦克利维 杰拉尔德 E. 平托 戴维 E. 朗克尔 著 劳兰珺 王祺 译
(Richard A. DeFusco) (Dennis W. Mcleavey) (Jerald E. Pinto) (David E. Runkle)

原书第2版



机械工业出版社
China Machine Press



QUANTITATIVE INVESTMENT
ANALYSIS (Second Edition)

定量投资分析

(美) 理查德 A. 德弗斯科 丹尼斯 W. 麦克利维 杰拉尔德 E. 平托 戴维 E. 朗克尔 著
(Richard A. DeFusco) (Dennis W. Mcleavey) (Jerald E. Pinto) (David E. Runkle)

劳兰珺 王祺 译

原书第2版



机械工业出版社
China Machine Press

Richard A. DeFusco, Dennis W. Mcleavey, Jerald E. Pinto, David E. Runkle. Quantitative Investment Analysis, 2nd Edition.

Copyright © 2004, 2007 by CFA Institute.

This translation published under license. Simplified Chinese translation copyright © 2012 by China Machine Press.

No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or any information storage and retrieval system, without permission, in writing, from the publisher.

All rights reserved.

本书中文简体字版由 John Wiley & Sons 公司授权机械工业出版社在全球独家出版发行。

未经出版者书面许可,不得以任何方式抄袭、复制或节录本书中的任何部分。

本书封底贴有 John Wiley & Sons 公司防伪标签,无标签者不得销售。

封底无防伪标均为盗版

版权所有,侵权必究

本书法律顾问 北京市展达律师事务所

本书版权登记号: 图字: 01-2011-3938

图书在版编目 (CIP) 数据

定量投资分析 (原书第 2 版) / (美) 德弗斯科 (DeFusco, R. A.) 等著; 劳兰珺, 王祺译. —北京: 机械工业出版社, 2012. 6

(CFA 协会投资系列)

书名原文: Quantitative Investment Analysis

ISBN 978-7-111-38802-9

I. 定… II. ①德… ②劳… ③王… III. 投资分析—定量分析 IV. F830.593

中国版本图书馆 CIP 数据核字 (2012) 第 126287 号

机械工业出版社 (北京市西城区百万庄大街 22 号 邮政编码 100037)

责任编辑: 蒋桂霞 版式设计: 刘永青

中国电影出版社印刷厂印刷

2012 年 7 月第 1 版第 1 次印刷

185mm × 260mm · 28 印张

标准书号: ISBN 978-7-111-38802-9

定价: 99.00 元

凡购本书, 如有缺页、倒页、脱页, 由本社发行部调换

客服热线: (010) 68995261; 88361066

购书热线: (010) 68326294; 88379649; 68995259

投稿热线: (010) 88379007

读者信箱: hzjg@hzbook.com



顾问委员会

(按姓氏拼音排序)

- 陈信华 上海大学经济学院金融系主任 教授
- 劳兰珺 复旦大学管理学院财务金融系 教授
- 李 翔 上海大学经济学院金融系 博士, CFA
- 李曙光 中国人民银行金融信息化研究所副所长
- 陆 军 中山大学岭南学院副院长 教授
- 孟庆轩 斯坦福大学研究员, 中国投资咨询公司首席经济学家
- 汪昌云 中国人民大学财政金融学院金融系主任 教授
- 吴冲锋 上海交通大学安泰经济与管理学院副院长 教授
- 叶永刚 武汉大学经济与管理学院副院长 教授
- 郑振龙 厦门大学研究生院副院长 教授
- 朱武祥 清华大学经济管理学院 教授

本书献给玛尔戈、雷切尔和丽贝卡

(理查德 A. 德弗斯科)

献给简、克莉丝汀和安迪

(丹尼斯 W. 麦克利维)

纪念特许金融分析师伊文 T. 范德胡夫

(杰拉尔德 E. 平托)

献给帕特丽夏、安妮和莎拉

(戴维 E. 朗克尔)

致中国读者

特许金融分析师（CFA）协会是一个由专业投资人士组成的全球性组织，并设定从业水平的优异标准。许多年来，我们已经成为金融市场职业道德行为的拥护者和全球投资领域受人尊敬的知识来源。我们最终的目标是通过创造把投资人的利益置于首位、市场处于最优运行状态、经济不断增长的环境，来做出更大的贡献。

在中国，CFA 协会发现人们对 CFA 称号有着很强的兴趣，这一称号被认为是全球最受尊重且最具投资可信度的。获得 CFA 称号需要参加 CFA 项目的学习并通过三个级别的考试，并成为 CFA 协会的正式会员。根据 2011 年 12 月和 2012 年 6 月的考试情况，中国的 CFA 注册考生人数已超过 30 000 人，CFA 中国（CFA 协会代表处）的会员也从 2009 年的 1 619 人增加到现在的 2 081 人^①。CFA 协会非常高兴看到这样的发展并将继续致力于帮助中国投资行业的发展。

CFA 协会投资系列适用于从研究生层次的金融专业学生到实务职业投资人的广泛的个人读者。针对中国市场出版的本系列中文版图书，标志着我们为满足不同全球投资市场不断增长的需求所付出努力的一个重要里程碑。本系列中的每本书都以清晰的案例讲解的方式，涵盖了投资工具、投资组合管理、财富管理、估值与分析技术，以及目前错综复杂的行业所需要的其他关键的主流和前沿概念。

我们希望这一版本能够提升你对金融行业的理解，并帮助你在当今快节奏的金融世界取得成功。

艾博科（Ashvin Vibhakar）博士，CFA

CFA 协会亚太区董事总经理

^① 根据 CFA 协会数据，2012 年 6 月 7 日。

丛书序

CFA 协会非常荣幸能够为你提供涉及投资领域专业内容的系列丛书，这些内容完全以对 CFA 项目的考生知识体系（CBOK）的高度重视为基础。CBOK 吸引了数以百计在投资操作方面的专家，成为了 CFA 三个级别考试的标杆。在本系列丛书发售的 2007 年，超过 120 000 名有抱负的投资从业人员投入逾 250 小时研读和学习这份材料及其他 CFA 项目的 CBOK 所提供的材料，以期获得梦寐以求的 CFA 特许证书。我们本着与 40 年来为投资从业人员颁发特许证书同样的理由提供这些材料，那就是提升资本市场服务人士的能力与道德水平。

来源

本系列丛书的一大重要特质即是渊源深厚。19 世纪 40 年代，一小部分人因共同的兴趣和工作发展成了社会团体，这就是我们如今认识的投资行业。

购买股票是一种投资行为而非赌场投机，人们形成这一观点至多不过几十年。在强盗贵族和股票市场大崩盘的插曲之后，为了规范操作市场而产生的美国证监会和相关法律距今也不过 10 年。

1945 年 1 月，一名来自哥伦比亚大学和格雷厄姆—纽曼公司的重要领军人物在 CFA 协会的《金融分析师杂志》中撰文指出，从事投资组合研究和管理工作的人士应当有一些凭证用以证明其能力和道德操守。这个人就是本杰明·格雷厄姆，股票分析之父、现代著名投资家巴菲特的老师。

这个关于创设相关证明的想法付诸实践仅仅用了 16 年。1963 年，284 名年逾 45 岁的勇士参与了考试并获得了相应的 CFA 证书。许多人并不十分了解的是，这一努力是根植于创造由专业从业人员为有需要的个人提供投资服务行业的渴求。这样才能创造出更加公平、有效的资本市场。

不论是医药、法律或是其他任何行业，都有其特定的品质特征。这些特质即是吸引严肃的人将毕生精力献给投资实践的原动力。第一，也是与这套书紧密相关的，行业应当具备知识体系；第二，行业会设立一定的准入要求，比如 CFA 特许证书；第三，从业者必须不断学习；第四，从

业者在服务时不可以获得个人利益为目的，也就是说，他应当合理处理个人事务并始终将客户利益放于首位，以公正之心在高效运作的全球资本市场进行专业投资。这将会鼓励人们将辛苦所得的存款本着合理、富有成效的目标进行再配置。

正如 CFA 创设时的执行官斯图尔特 C. 谢泼德所说：“社会要求一种行业及其从业者不仅仅是贸易、艺术和商业技工，其必须公开担保已经建立并仍维持准入资格、遵循公平交易原则和操作系统、持续学习以服务客户，而这也将为从业者赢得地位、声誉和独立的决定权。他们被期望很多，但经受住时间的考验，收益将是丰厚的。”

由美国心理协会、美国教育研究协会和全美教育测评委员会提出的“教育和心理测试标准”，证明了专业证书考试的正确性应当主要通过核查考试内容是否正确体现专业的操作来予以说明。与此同时，实践分析案例应当成为建立其考察可靠性的基础，因为这可以确定被考核者是否具备了从业者应有的知识与技能。

40 多年来，成千上万服务于 CFA 协会课程委员会的从业者和学者从所有与投资相关的概念和观点中层层筛选，以求建立一套知识体系，设立 CFA 的相关教程。CFA 协会课程发展的一大标志是其得到了各阶段从业者的广泛应用。

CFA 协会自 1995 年起始终努力遵循一套有条理的实践分析流程，包括召开特殊实践分析论坛，这些论坛已在全世界 20 多个地区展开，得到的结果交给 70 000 名 CFA 特许分析师，由他们确认核查其衍生的知识体系。

对读者而言，这意味着本书中包含的概念是由此领域实际操作中的专家总结得来，他们深刻明白在这一行业中成功的从业者所需承担的责任和具备的知识。我们非常愿意为本系列丛书的读者付出种种努力。

优势

本系列丛书的广泛适用性体现在，它不仅适合刚刚接触资本市场、努力尝试进入这一竞争极其激烈的投资管理领域的新手，同样适合试图寻找到更新自身知识便捷方式的专业老手。本书所有的章节都为想要深入探究特定概念的读者提供了丰富的参考文献。

对于新手来说，本书提供的任何投资专家都必须掌握的基本概念，配合大学学习、财经杂志阅读，将会帮助你更好地理解投资行业。我们坚信大众的确低估了成为最杰出的公司及个人所需接受的训练。本书中给出了许多成功公司训练时使用的基本内容。如果没有对这方面的基本了解和对资本市场如何运作来给股票合理定价的理解，你也许无法成为一名强有力的竞争者。不仅如此，这些概念能带给你对此类工作的真实感受，即日复一日地管理投资组合、进行研究和相关的努力。

投资行业尽管相对而言报酬非常高，却并非人皆可为。这份工作需要你有一种特质，能够从根本上理解和吸收书本中所教授的知识，并在现实操作中将其转化应用。实际上，进入这一行业的大多数人无法长期从事该职业。有抱负的从业者应当给自己一段时间认真思考自己是否适合这一行业。要做出这样一个至关重要的决定，除了以阅读和考量这一行业的信条来让自己做好准备以外别无他法。

较有经验的从业者明白，资本市场的特性决定了必须坚持不断学习。市场总在快速发展，同时聪明的头脑总是能找到新方法来制造敞口、吸引资本或是控制风险。许多本书上提到的概念在10年甚至20年前我们开始从事此行业时尚未出现。对冲基金、衍生品、另类投资概念和行为金融学都是新应用和新概念的例子，在最近几年它们已经改变了资本市场。随着市场自身的创造与再创造，一套质量上乘的基础性投资丛书是极有价值的。

我们中那些在这行业中摸爬滚打过的从业者已然知晓，我们必须不断地打磨自己的技能和知识，才可能与不断出现的新秀们竞争。实际上，当我们与一些主要的雇员就其培训意愿交流时，我们总是被告知，他们面对的最大挑战是如何帮助资深专家在极其紧张的时间和工作压力下，能够跟上专业潮流并与不断更新知识的同行保持联系。本系列丛书将提供对这个问题的部分回答。

传统观点

对于机敏的从业者而言，他们很快就会发现市场的两大特性。第一，价格是根据传统观点或者一个含有许多市场变量的函数确定的。实质上，在市场中，真理就是市场的观点，以及市场如何据此确定借贷定价。第二，因为传统观点是普遍理论和学习的发展成果，它当然常有瑕疵或者至少会随外界条件改变而变得不再合理。

当我于20世纪70年代中期初入此领域时，传统观点认为书中考察的概念对于这个竞争激烈的市场而言太过学术化了。许多人认为当时最好的投资公司是由信奉中庸思想、凭借市场敏锐直觉和有良好从业记录的人领导的。在这个从业者良莠不齐的世界里，一些概念曾被认为是无用的。传统观点还能错得更加离谱吗？若是，那我真不知道会是几时。

在我就职的这些年，发展顺利的专业投资管理公司都充满了坚定、智慧、好奇的投资专家，他们努力将这些概念以一种严谨且训练有素的方式加以应用。如今，最好的公司是由这样一些人来带领，他们谨慎地做出投资预测并通过市场严密地验证它们——是否按照定量策略得出，是否是行业内比价购买的结果，是否是对冲基金的策略产物。他们的目标是创造能够依统计数据的可靠性复制的投资流程。我相信信奉所谓的学术派不会比现实世界中的投资管理者更加成功。

关于本系列丛书

本系列最初的四本书阐述了近 35% 的考生知识体系。新增的关于公司理财和国际财务说明分析的部分仍在不断发展中，而且更多的主题即将出现。

其中一本书这些年来在投资管理行业颇为突出，它是马金和塔特尔的《投资组合管理：动态过程》。第 3 版在 1990 年第 2 版的基础上更新了主要概念。许多协会内的资深成员（比如我）已经有了前两版，也准备将这第 3 版收入囊中。该书不仅吸收了其他文献中的概念，并将它们置于投资组合这一主题之下，同时还更新了一系列概念，例如，另类投资、业绩报告标准、投资组合操作和非常重要的个人投资者投资组合的管理。长期以来我们都关注着机构投资者的投资组合，引起读者对个人投资者投资组合的注意，将会成为此版较之前两个版本的重大改进。

《定量投资分析》着眼于当今专业投资者必需的一些核心工具。该书不仅包含经典的货币时间价值理论、现金流贴现应用、或然性内容，还包含超越了传统思考方式的非常有价值的两方面。

第一，为了检验假设，一些章节在公式中运用了相关和回归。这考验了许多从业者的一项重要能力：去粗取精。对于大多数的投资研究员和管理者而言，他们的分析不仅仅是新近产生数据的结果，还包括了对业绩优劣的考量。当然，他们也会综合和分析其他人的主要研究报告。如果不以严苛的态度去理解质量报告，那你不仅不能很好地研究，甚至将无法胜任要求不那么严格的研究。现实世界中从不缺乏徒有数量分析而全无正确性和严密性的研究。

第二，关于投资组合概念的最后一章将超越传统资本资产定价模型（CAMP）类的工具，带领读者进入由多因素模型和套利定价理论组成的更真实世界。过去，许多人已经感觉到了 CAMP 与实际操作中的鸿沟，这一章将会助你解开困惑。

《股权资产估值》对于从事证券估值、需要理解证券定价的读者必是一本令人信服且极其重要的读物。见多识广的从业者知道证券评估的几种常见形式——股利贴现模型、自由现金流模型、市盈率模型和剩余收益模型（通常都以交易名称为人们所知），都能够在特定的假设前提下互换使用。对于这些根本假设有了较为深刻的理解以后，专业投资者就能更好地理解其他投资者在何种假设之下计算其估值。在我先前担任证券投资团队领导者的经历中，这部分知识为我们与其他投资者的竞争中赢得了优势。

《固定收益证券分析》的内容是近几年才出现的前沿概念，并且已经在过去极大地扩展了范围。对并非老练的从业者来说，这是最新的一部分材料，他们往往不是固定收益方面的专家。对期权和衍生品在一度保守的固定收益知识的应用使固定收益分析在这个领域兴盛起来。这不仅挑战从业者，逼着他们熬夜学习以赶上信用衍生产品、互换期权、担保抵押证券、按揭证券和其他

产品的发展，同时还给全球的中央银行施加了监管并提防相关事件风险的巨大压力。通过对这些新发现的彻底认识，专业投资者能够更好地估计和理解中央银行家和市场面临的挑战。

无论你是刚刚上路的新手，抑或鬓发斑白的老将，我都希望这一系列全新的丛书能够在增长投资知识的道路上助你一臂之力，使你能够紧紧跟上日新月异的市场环境。作为一个长期坚定服务于专业投资者的非营利性协会，CFA 协会非常荣幸能够为你提供这样的机会。

杰夫·狄尔梅尔 (Jeff Diermeier), CFA

CFA 协会总裁兼首席执行官

2006 年 9 月

前言

定量投资分析如何优化投资组合决策

我是一名数量分析专家。就我自身的经历来说，我会在管理投资组合时使用定量投资技术。然而，每当我对别人说我是一名数量分析专家时，他们常常会反问道：“但是马克，你不是一直是律师吗？”呃，确实是这样的，但是……

事实上，数量分析专家来自各行各业。无论我们被称为数量分析专家（quants）、股市分析高手（quant jocks）、计算机程式设计者（gear heads）、计算机专家（computer monkeys），还是任何其他用来形容那些喜欢在一堆纸上潦草地书写方程式的投资者的绰号，我们都具有一个共同的特征——使用定量分析方法来做出更好的投资决策。你不必在艰深难懂的数学领域成为一名具有博士学位的火箭专家（rocket scientist）之后，才能成为一名合格的数量分析专家（虽然我认为这样的情况是存在的。许多以前是火箭专家的人士由于发现在金融市场工作非常有趣，而且获益丰厚，从而成了数量分析专家）。任何人都可以成为一名数量分析专家，甚至是一位律师。

但是，让我们退一步来思考一下。为什么投资者会想要在投资组合管理中使用定量分析工具呢？数量分析专家如此受欢迎的原因主要有如下三点。

首先，金融市场是一个非常复杂的场所。在市场中存在着许多交织在一起的变量，这些变量都会影响到一个投资组合中的证券价格。例如，上市公司的股价可能会受到诸如利率、经常账户赤字、政府支出以及经济周期等宏观因素的影响。这些因素可能会影响一家公司为新项目融资所采用的资本成本，或者会影响一家公司客户的消费方式，或者会通过政府支出项目为经济提供增长动力。

除了宏观变量之外，一家公司股票的价值也会受到公司自身特殊因素的影响。诸如现金流量、运营资本、账面市值比率、盈利增长率、股利政策以及负债权益比率等因素也会对每家上市公司的价值产生影响。这些因素被认为是对于特定公司而不是对于广泛的股票市场产生影响的基本面因素。

其次，我们还可以考察影响公司估值的金融市场变量。公司的“贝塔”或者系统风险的度量

指标将会影响公司的预期收益率，并最终影响到该公司的股价。著名的度量股票贝塔的资本资产定价模型只是本书第8章中所讨论的各种类型模型中的一种线性回归方程式。

最后，行为变量也会影响证券的价值。诸如羊群效应、过度自信、对于盈利公告的过度反应以及惯性交易等行为都会影响公司股票的价格。这些行为变量会对股票价格产生一个持续性的影响（还记得1998~2001年的互联网泡沫吗？当时高科技股票似乎要充斥整个世界），同时它们还会在股票真实价值周围产生大量的“噪声”（noise）。

如果没有分析各个变量对于证券价格影响的具体框架，那么通过一次性地对于所有这些变量加以考察，以确定出某个证券的真实价值，可能是一项巨大的工程。仅凭人类的头脑（至少，仅凭我的头脑）显然是不可能有能力以一种严谨的方式衡量出各家公司的特有因素（如市盈率）、宏观经济变量（如政府支出项目）、投资者行为模式（如惯性交易）以及其他可能的潜在变量所造成的影响的。

这就是《定量投资分析》一书可以给予我们帮助的地方。像第11章中所介绍的因素建模技术，可以用来作为人类直觉判断的一个补充，并由此得出一个融合大量可能对证券价格产生影响的变量的定量分析框架。进一步地，给定可能会对证券价值造成影响的许多变量，那么我们就不能孤立地来考虑每个变量。引起某个证券价格上升或者下降的经济因素是以一种复杂的网络方式交织在一起的，因此，我们必须将变量放在一起加以考虑，以确定出它们对于该证券价格的共同影响。

这就是第8章和第9章最有价值的地方。这两章提供了我们建立回归方程来研究对证券价格产生影响的经济因素的基本概念。第8章和第9章中所提供的回归技术可以用来筛选出哪些是对证券价格产生显著影响的变量，哪些仅仅是提供“噪声”的变量。

另外，第9章为读者介绍了“虚拟变量”（dummy variables）的概念。尽管它们的名称中含有木偶的意思，但是我们不必像一个木偶一样使用虚拟变量。虚拟变量是研究市场的不同状态以及对于证券价格影响的一种灵活方式。它们常被称为“二元”（binary）变量，因为它们将市场划分为两种状态的观测值。例如，金融市场上升阶段与金融市场下降阶段；共和党在白宫执政阶段与民主党在白宫执政阶段；芝加哥小熊队获胜（几乎没有过）或者芝加哥小熊队失利……诸如此类。上面列举的最后一个变量——芝加哥小熊队的战绩，我可以确信它对于证券估值没有影响，虽然我是一名长期支持并承受其失利痛苦的小熊队的球迷，它确实给我的情绪带来了影响。

作为另外一个例子，最近我写的一份研究报告中，我对于私募基金经理的行为进行了研究，研究的内容是私募基金经理对于其私募基金投资组合的定价是否依赖于公开股票市场表现的好坏。为了进行这项研究，我使用虚拟变量将市场分成了两个状态：上升的公开股票市场以及下降的公开股票市场，并在这些虚拟变量的基础上进行了回归。通过这种方式，我观测到了私募基金经理如何根据公开股票市场的表现来调高或者调低他们的私募基金投资组合估值的不同行为

模式。

《定量投资分析》一书将会给读者带来价值的第二个原因是，它提供给我们考察广泛的经济因素和证券的基本工具。这不仅仅因为存在许多影响证券价值的互相交织的经济因素，而且因为市场中大量的证券可能会令人望而却步。因此，大部分投资者只会关注市场中那些可投资证券中的一小部分。

让我们考虑一下美国股票市场。一般来说，该市场中的股票可以根据公司规模分为三类：大市值、中市值和小市值股票。这种分类很少是因为在估值中可能存在“规模”效应，更多的是因为考虑到实际应用上的限制，即资产管理经理无法分析超出一定数量的股票。因此，传统的基本面投资者会选择美国股票市场的不同部分来进行他们的证券分析。然而，将股票市场划分为不同的规模类别，实际上为投资经理设置了障碍。例如，我们无法解释为什么一个了解盈利冲击是如何影响股票价格的投资组合经理不能在包含所有市值股票的市场中进行投资。

这就是第6章和第7章可能对我们有用的地方。抽样、估计和假设检验这些定量技术都可以用来对大量数据进行分析。这使得投资组合经理可以打破传统的诸如按市值分类投资的限制，从而在广泛的股票集合中进行投资。从这个角度来看，定量分析并没有取代传统资产管理经理所采用的基本面选股技术。相反，定量分析将投资组合经理对于公司、宏观经济以及市场变量的了解拓展到了一个更加广泛的投资机会集之中。

第3章和第4章所给出的统计工具和概率论概念同样也会给我们带来帮助。所要分析的数据集的规模越大，则由该数据集所得到的参数估计的可信度就越高。经济分析的广度不仅会改善定量分析在统计上的可信度，而且会增加经济因素和股票价格变动之间关系的可预测性。本书中所提供的统计工具能帮助投资组合经理在更大的证券投资集中（而不是在之前的小范围中）发挥出其全部的分析技术能力。

另外一个例子也可以说明这个问题。加利福尼亚公共雇员退休系统（California Public Employees' Retirement System, CalPERS）每年会公布一份美国治理最差公司的名单。该名单已经连续发布了16年，并且取得了非常大的成功（非常受人们欢迎）。在发布该名单的早期过程中，我们是从美国股票市场的一个子集中进行样本选取的。然而，之后该过程发展到了考察CalPERS投资组合中持有的每一只美国股票，而不管该股票的市值是多少。这样做的结果就是要求我们每年基于各种经济因素和公司治理变量分析多达1800只股票。因此，我们如果没有应用定量筛选工具来扩大CalPERS投资组合的治理情况分析范围的话，我们几乎不可能对该数据样本中的大量证券进行分析。

最后，《定量投资分析》一书还可以为我们提供投资过程中的一些规则。我们都是人，作为一个人，我们都会犯错误。如果我在这里详细叙述我在职业生涯中所有不断重复犯的投资错误的话，那么这篇前言将超过本书所有章节的长度之和。举一个简单的例子，星巴克咖啡店是我比较

喜欢去的场所之一，在某天早上星巴克刚开门的时候，我步入了它的一家连锁店，当时我想了解一下里面的嘈杂声到底是怎么回事。在那时，我发现拿铁咖啡售价仅为 1.50 美元。我记得当时我想这真是一个令人难以置信的价格，并且我非常清晰地记得，我自言自语道：“噢，太不可思议了，这简直无法令人理解。”

因此，让我们回到定量分析技术的讨论之中，这些技术是如何帮助我们的呢？在上面这个例子中，这些技术就可以帮助我们排除人类思想中的偏见，并更加具有逻辑性地对星巴克的前景进行思考。如果我花点时间使用本文中所提供的定量工具进行实证研究的话，我将会了解到该 1.50 美元拿铁的真实价值到底是多少。

现实中的真实情况是，我们都受到诸如过度自信、惯性以及过度反应等行为偏见的影响。我们不仅需要对于上面讨论的这些内容进行分析，而且需要在我们做出投资决策时，解释并体现它们。对于投资者来说，可能需要克服的最大行为障碍是无法在证券价格下跌时果断地将其卖出。通常的情况是我们会对投资组合中的证券倾注过多的感情，以致于发现自己难以在价格开始下跌时果断卖出它们。

但是，这些正是定量投资分析会帮助到我们的地方，因为它是不会受到感情左右的。定量工具和建模技术可以从投资组合决策过程中排除感情和认知上的偏差。作为投资组合经理，我们需要的是客观的评价、批判性的思维以及严格的标准。不幸的是，我们平时所养成的习惯和思维方法有时会闯入进来。然而，定量分析模型是不带感情的，它们能够排除我们在认知上的偏见，这种偏见在某种程度上是由于我们无法通过观察镜子认识到真实的自己而产生的。（事实上，当我在照镜子时，我所看到的是一个身高 1.93 米的帅小伙，但是之后我的妻子玛丽会提醒我说，其实我只有 1.85 米，并且她曾经拥有比我长得更帅的追求者。）

总而言之，投资者将会从本书所提供的方法、模型和技术中获益良多。对于那些刚开始在投资组合管理过程中使用定量工具的投资者来说，本书是一本很好的入门书籍。同样地，对于那些已经认识到定量分析重要性的人士来说，本书也是一本极佳的参考指导手册。定量投资并不难掌握，甚至像我这样一名律师也可以做到。

马克 J. P. 安森 (Mark J. P. Anson)

赫尔墨斯养老管理基金 (Hermes Pensions Management) CEO

英国电信养老金计划 (British Telecomm Pension Scheme) CEO

mark@hermes.co.uk

致 谢

我们要感谢在本书出版过程中发挥重要作用的许多人士。

罗伯特 R. 约翰逊 (Robert R. Johnson), CFA, CFA 协会和 CIPM 项目部常务董事。正是他看到了这门专业课程在教学资料上的需求,从而发起了这项书籍出版计划。我们非常感谢他对于本教材多次修改的支持。CFA 项目部的高级主管们在本书两个版本的写作过程中无私地花费宝贵的时间给予我们许多建议。菲利普 J. 杨 (Philip J. Young), CFA, 在本书学习成果报告部分的写作过程中曾多次提供帮助,并参加了最终手稿的审阅。简 R. 斯克维尔斯 (Jan R. Squires), CFA, 强调了本书的写作目的并提出了可尝试的写作方向。玛丽 K. 埃里克森 (Mary K. Erickson), CFA, 对于全书的准确性做出了贡献。约翰 D. 斯托 (John D. Stowe) 对于一些章节的修改提出了建议。

候选人课程委员会的管理顾问委员会为我们提供了有力的人员支持:詹姆斯·布朗森 (James Bronson), CFA, 主席;彼得·麦基 (Peter Mackey), CFA, 刚卸任的前主席;以及其成员,阿兰·梅德尔 (Alan Meder), CFA, 维多利亚·拉缇 (Victoria Rati), CFA, 马特·斯坎兰 (Matt Scanlan), CFA, 以及 CFA 课程委员会办事机构。

对于该版本原稿的审阅者包括:小菲利普·法娜拉 (Philip Fanara, Jr.), CFA;简·法里斯 (Jane Farris), CFA;戴维 M. 杰索普 (David M. Jessop);丽萨 M. 乔布兰科 (Lisa M. Joubanc), CFA;埃斯杰特 S. 兰芭 (Asjeet S. Lamba), CFA;玛丽奥·拉瓦利 (Mario Lavallee), CFA;威廉 L. 伦道夫 (William L. Randolph), CFA;埃里克 N. 雷默尔 (Eric N. Remole);维贾伊·辛盖尔 (Vijay Singal), CFA;佐伊 L. 范·辛德尔 (Zoe L. Van Schyndel), CFA;夏洛特·威姆斯 (Charlotte Weems), CFA;以及拉翁 F. 惠特梅尔 (Lavone F. Whitmer), CFA。我们感谢他们所做出的杰出工作。

我们同样要感谢那些使用过本书第一版并给我们提出许多意见的读者。

雅克 R. 盖聂 (Jacques R. Gagne), CFA, 格里高利 M. 诺隆哈 (Gregory M. Noronha), CFA, 以及桑吉夫·萨巴沃尔 (Sanjiv Sabherwal) 非常仔细地校对了每个章节。我们对他们忘我的辛勤工作表示感谢。我们也要感谢萨巴沃尔博士在修改各章中的例子以及编写各章末的一些问题及解

答方面给予我们的专业指导与帮助。

菲奥娜 D. 罗素 (Fiona D. Russell) 对本书进行了全面的文字校对工作, 在本书的准确性和可读性方面做出了巨大的贡献。

CFA 项目部负责修改的项目经理旺达 A. 劳泽尔 (Wanda A. Lauziere) 对于本书原稿从策划到出版的整个过程予以了非常专业的指导, 并且在各个方面的修改过程中提供了许多帮助。

最后, 我们要感谢芝加哥伊博森协会为我们无私地提供了 EnCorr Analyzer™ 软件。

目 录

致中国读者 (艾博科)	
丛书序 (杰夫·狄尔梅尔)	
前言 (马克 J. P. 安森)	
致谢	
第 1 章 货币的时间价值	1
1.1 引言	1
1.2 利率:经济学的解释	1
1.3 单笔现金流的将来值	3
1.3.1 复利的频数	7
1.3.2 连续复利	8
1.3.3 报价利率和有效利率	9
1.4 现金流序列的将来值	10
1.4.1 等额现金流序列——普通年金	10
1.4.2 不等额现金流序列	12
1.5 单笔现金流的现值	12
1.5.1 求解单笔现金流的现值	12
1.5.2 复利的频数	14
1.6 现金流序列的现值	15
1.6.1 等额现金流序列的现值	15
1.6.2 无限期等额现金流序列的现值——永续年金	18
1.6.3 始点不在零时刻的现金流序列的现值	19
1.6.4 不等额现金流序列的现值	20
1.7 求解利率、期数或年金支付额	21
1.7.1 求解利率和增长率	21
1.7.2 求解期数	23
1.7.3 求解年金支付额	24
1.7.4 现值和将来值换算关系的回顾	27
1.7.5 现金流可加性原理	28
第 2 章 贴现现金流的应用	30
2.1 引言	30
2.2 净现值和内部收益率	30
2.2.1 净现值和净现值准则	31
2.2.2 内部收益率和内部收益率准则	32
2.2.3 与内部收益率准则相关的问题	35
2.3 投资组合收益的度量	37
2.3.1 货币加权收益率	37
2.3.2 时间加权收益率	38
2.4 货币市场收益率	42
第 3 章 统计学概念和市场收益率	47
3.1 引言	47
3.2 一些基本概念	47
3.2.1 统计学的本质	48

3.2.2	总体和样本	48
3.2.3	度量尺度	49
3.3	用频数分布汇总数据	50
3.4	数据的图形表示	56
3.4.1	直方图	56
3.4.2	频数多边形和累积频数 分布图	58
3.5	集中趋势的度量	59
3.5.1	算术平均数	60
3.5.2	中位数	63
3.5.3	众数	65
3.5.4	有关均值的其他概念	66
3.6	位置的度量:分位数	73
3.6.1	四分位数、五分位数、十分 位数、百分位数	73
3.6.2	分位数在投资中的应用	77
3.7	离散度的度量	78
3.7.1	极差	79
3.7.2	平均绝对偏差	79
3.7.3	总体方差和总体标准差	81
3.7.4	样本方差和样本标准差	83
3.7.5	半方差、半离差及其相关 概念	86
3.7.6	切比雪夫不等式	87
3.7.7	变异系数	88
3.7.8	夏普比率	90
3.8	收益率分布的对称性和偏度	92
3.9	收益率分布的峰度	96
3.10	使用几何平均和算术平均	100

第4章	概率论中的一些概念	102
4.1	引言	102
4.2	概率、期望值和方差	102

4.3	投资组合的期望收益和收益的 方差	120
4.4	概率论的一些议题	126
4.4.1	贝叶斯公式	126
4.4.2	计数原理	130

第5章	常用概率分布	134
5.1	引言	134
5.2	离散型随机变量	134
5.2.1	离散均匀分布	136
5.2.2	二项分布	137
5.3	连续型随机变量	145
5.3.1	连续均匀分布	145
5.3.2	正态分布	147
5.3.3	正态分布的应用	153
5.3.4	对数正态分布	155
5.4	蒙特卡罗模拟	159

第6章	抽样和估计	165
6.1	引言	165
6.2	抽样	165
6.2.1	简单随机抽样	166
6.2.2	分层随机抽样	167
6.2.3	时间序列数据和横截面数据	168
6.3	样本均值的分布	170
6.4	总体均值的点估计和区间估计	173
6.4.1	点估计量	173
6.4.2	总体均值的置信区间	174
6.4.3	样本量的选择	179
6.5	抽样中的若干问题	180
6.5.1	数据挖掘的偏差	180
6.5.2	样本选择的偏差	183
6.5.3	前视偏差	184
6.5.4	时期偏差	184

第7章 假设检验 186

7.1 引言 186

7.2 假设检验 187

7.3 关于均值的假设检验 195

7.3.1 对单个均值的检验 195

7.3.2 对均值间差异的检验 201

7.3.3 对（配对样本）均值差的检验 204

7.4 关于方差的假设检验 207

7.4.1 对单个方差的检验 207

7.4.2 对两个方差是否相等的检验 209

7.5 其他议题：非参数推断 211

7.5.1 相关性检验：斯皮尔曼（Spearman）秩相关系数 212

7.5.2 非参数推断：总结 214

第8章 相关性和回归 215

8.1 引言 215

8.2 相关性分析 215

8.2.1 散点图 215

8.2.2 相关性分析 216

8.2.3 计算和解释相关系数 218

8.2.4 相关性分析的局限 220

8.2.5 相关性分析的应用 222

8.2.6 相关系数显著性检验 227

8.3 线性回归 230

8.3.1 单变量的线性回归 230

8.3.2 线性回归模型的前提假设 ... 232

8.3.3 估计量的标准误差 234

8.3.4 决定系数 236

8.3.5 假设检验 238

8.3.6 单变量回归中的方差分析 ... 243

8.3.7 预测区间 246

8.3.8 回归分析的局限 248

第9章 多元回归和回归分析中的问题 249

9.1 引言 249

9.2 多元线性回归 249

9.2.1 多元线性回归模型的前提假设 254

9.2.2 预测多元线性回归模型中的因变量 258

9.2.3 检验是否所有回归系数为零 ... 258

9.2.4 调整后的 R^2 平方 260

9.3 虚拟变量在回归中的使用 261

9.4 回归假设的违背 264

9.4.1 异方差 265

9.4.2 序列相关 269

9.4.3 多重共线性 273

9.4.4 异方差、序列相关、多重共线性：问题的总结 275

9.5 模型设定和设定中的错误 276

9.5.1 模型设定的原则 276

9.5.2 函数形式误设定 277

9.5.3 时间序列误设定（自变量与误差相关） 283

9.5.4 其他类型时间序列误设定 ... 285

9.6 因变量是定性变量的模型 286

第10章 时间序列分析 288

10.1 引言 288

10.2 处理时间序列数据所面临的挑战 289

10.3 趋势模型 290

10.3.1 线性趋势模型 290

10.3.2 对数线性趋势模型 292

10.3.3 趋势模型和误差项相关性	
检验	296
10.4 自回归时间序列模型	297
10.4.1 协方差平稳序列	297
10.4.2 检测自回归模型中的序列	
相关误差	298
10.4.3 均值回复	301
10.4.4 多期预测和预测的链式	
法则	301
10.4.5 比较预测模型的表现	304
10.4.6 回归系数的不稳定性	306
10.5 随机游走和单位根	308
10.5.1 随机游走	308
10.5.2 非平稳数据的单位根	
检验	311
10.6 移动平均时间序列模型	314
10.6.1 用 n 期移动平均平滑历史	
数据	314
10.6.2 用移动平均时间序列模型来	
进行预测	316
10.7 时间序列模型中的季节性	317
10.8 自回归移动平均模型	321
10.9 自回归条件异方差模型	322
10.10 两个以上时间序列的回归	324
10.11 时间序列的其他议题	328
10.12 时间序列预测建议采取的步骤	328
第11章 投资组合的概念	331
11.1 引言	331
11.2 均值方差分析	331

11.2.1 最小方差前沿及其相关	
概念	332
11.2.2 扩展到3种资产的情况	339
11.2.3 多个资产最小方差前沿的	
确定	341
11.2.4 分散化和投资组合的规模 ...	344
11.2.5 存在无风险资产条件下的	
投资组合选择	347
11.2.6 资本资产定价模型	353
11.2.7 均值方差投资组合选择	
规则：一个介绍	355
11.3 均值方差分析在应用中的	
问题	358
11.3.1 估计均值方差优化问题	
中的输入参数	358
11.3.2 最小方差前沿的不稳定性 ...	362
11.4 多因素模型	365
11.4.1 因素和多因素模型的类型	366
11.4.2 宏观经济因素模型的结构	367
11.4.3 套利定价理论和因素模型	369
11.4.4 基本面因素模型的结构	374
11.4.5 多因素模型在当前实践中的	
运用	374
11.4.6 应用	380
11.4.7 总结	392
附录	394
术语表	403
参考文献	416
CFA 项目介绍	422
作者简介	423
译者后记	424

货币的时间价值

1.1 引言

作为个人，我们经常会面临为未来之需存款，或者为当前消费借款的决策问题。这时，如果是进行存款，我们需要决定存款的金额，而如果是为购物而申请贷款，那么我们就需要决定借款的金额。作为投资分析师，我们的许多工作要涉及对各类交易进行评估，这些交易往往具有当前与未来的现金流。例如，当我们给一个证券定价时，我们就是在试图确定一系列未来现金流的值。要想使得上述所有的工作都能够正确地完成，我们就必须理解货币时间价值问题中的数量关系。货币具有时间价值的原因在于，对于给定数量的一笔钱，人们认为越早收到，其价值越高。因此，当前数额较小的一笔资金可能与未来获得的一笔数额较大的资金具有相同的价值。货币的时间价值（time value of money）作为量化投资分析中的一个议题，讨论的就是不同时点现金流之间的等价转换关系。掌握货币时间价值的概念和计算方法对于投资分析师是最基本的要求。

本章的内容安排如下：第2节介绍一些在本章中将会用到的术语，并为我们将要讨论的变量提供一些经济学上的直观解释；第3节着手解决当前的一笔投资在未来时点上的价值的确定问题；第4节阐述现金流序列将来值的确定问题。上述两节提供将单笔现金流或现金流序列换算成在未来时点上的对等价值的一些方法；第5节和第6节分别讨论单笔未来现金流和未来现金流序列在当前时点的换算值；在第7节中，我们探讨如何确定货币时间价值问题中其他相关的变量。

1.2 利率：经济学的解释

在本章中，我们将频繁地提到利率。在一些例子中，我们会假定利率是一个特定的数值；而在另一些例子中，利率将会是我们需要去求解的未知变量。在讨论货币时间价值问题的原理之前，我们必须阐述一些基本的经济学概念。在本节中，我们简要地给出利率的含义及解释。

货币的时间价值考虑的是发生在不同日期的现金流之间的换算关系。换算关系的想法相对比较简单。考虑如下交易：今天你支付10 000美元，同时作为回报收到9 500美元。你会愿意接受这样的安排吗？不大可能。但是如果你在今天收到9 500美元，而在1年之后才支付10 000美元，

你会接受这个交易吗？这两笔钱能够认为是等价的吗？有可能的，因为1年后支付的10 000美元的价值对你来讲可能低于今天支付的10 000美元的价值。因此，对在1年后收到的10 000美元进行贴现（discount）是公平的；也就是说，根据付款前所要经过的时间长短来削减货币的价值。利率（interest rate），记为 r ，是一种收益率，它反映了不同日期发生的现金流之间的关系。如果今天的9 500美元和1年之后的10 000美元在价值上是相等的，那么 $\$10\,000 - \$9\,500 = \$500$ 就是对于1年后而不是现在收到10 000美元所要求的补偿。利率表示为以收益率形式表示的补偿要求，为 $\$500/\$9\,500 = 0.0526$ 或者为5.26%。

我们可以从三个角度来认识利率。首先，利率可以被认为所要求的收益率，即投资者接受某项投资所要求获得的最低收益率。其次，利率可以视为贴现率。在上面的例子中，5.26%就是我们用来贴现未来的10 000美元以求得其当前价值的一个比率。因此，我们几乎总是交替使用着“利率”和“贴现率”这两个术语。再次，利率可以看做是机会成本。机会成本（opportunity cost）是指投资者由于采取某项特定的行动而放弃的价值。在上面的例子中，如果提供9 500美元的这一方选择直接在今天消费这笔钱而不是借给另一方，那么他就放弃了获得5.26%收益的机会，所以我们将5.26%视为当前消费的机会成本。

经济学告诉我们，利率是由市场中的供给和需求的力量所决定的，投资者是资金的供给方，而借款者是资金的需求方。如果从投资者的角度来分析由市场决定的利率，我们可以将利率 r 看做由实际无风险利率和4种风险溢价所组成的。所谓风险溢价就是指由于承担不同的风险所需要获得的额外报酬或补偿。

$$\begin{aligned}\text{利率} = & \text{实际无风险利率} + \text{通货膨胀风险溢价} + \text{违约风险溢价} \\ & + \text{流动性风险溢价} + \text{到期日风险溢价}\end{aligned}$$

- 实际无风险利率（real risk-free interest rate）是指在没有通货膨胀预期条件下，完全无风险证券的单期利率。在经济理论中，实际无风险利率反映的是人们选择在当前或未来进行实际消费的时间偏好。
- 通货膨胀风险溢价（inflation premium）是对投资者所面临的预期通胀的补偿，它反映了从当前到债务到期期间预期的平均通货膨胀率。通货膨胀降低了单位货币的购买力，即单位货币可以购买的商品或服务的数量。实际无风险利率与通货膨胀风险溢价之和称为名义无风险利率（nominal risk-free interest rate）^①。许多国家有政府短期债券，其利率可被认为代表了该国家的名义无风险利率。例如，美国90天的国库券利率代表了那个时间段美国的名义无风险利率^②。美国国库券可以以很低的交易成本大量地进行买卖，并且得到美国政府的履约承诺和信用的支持。
- 违约风险溢价（default risk premium）是由于存在借款人未按照合同约定的时间和数额履行所承诺的支付的可能性而给予投资者的补偿。
- 流动性风险溢价（liquidity premium）是由于在投资需要迅速变现的情况下，存在相对于

① 理论上来说，1加上名义利率等于1加上实际利率和1加上通货膨胀率的乘积。然而，作为一种简便的近似，名义利率就等于实际无风险利率和通货膨胀风险溢价之和。在此处的讨论中我们利用这一近似可加关系来重点阐述基本概念。

② 其他发达国家发行与美国国库券类似的证券。法国政府发行期限为3个月、6个月和12个月的BTFs或者可转让固定利率贴现国库券；日本政府发行期限为6个月和12个月的短期国库券；德国政府以贴现价格发行期限长达24个月的财政融资票据和财政贴现票据；英国政府发行期限从1天到364天的金边国库券；而加拿大的政府债券市场与美国市场密切相关，发行期限为3个月、6个月和12个月的加拿大国库券。

投资公允价值而言的损失风险，而给予投资者的补偿。例如，美国国库券不具有流动性风险溢价，因为人们可以在市场上大量地买卖而不影响其市场价格。相反，许多小公司发行的债券在发行后交易不频繁，因此，这些债券的利率中就包含有流动性风险溢价，以弥补投资者在出售时所面临的较高成本（包括对于价格的冲击）。

- 到期日风险溢价（maturity premium）是指通常（在其他条件相同的情况下）随着到期日的延长，债券的市场价值对于市场利率变动的敏感性增加，为此而给予投资者的补偿。较长期限、流动性好的国债和短期国债之间的利差反映了长期债务具有正的到期日风险溢价（以及可能的不同的通货膨胀风险溢价）。

利用以上关于利率经济含义的认识，我们现在可着手讨论解决货币时间价值的问题。我们首先从单笔现金流的将来值问题开始讨论。

1.3 单笔现金流的将来值

在本节中，我们介绍与单笔现金流或者一次性投资相关的时间价值问题。我们将阐述初始投资或者现值（present value, PV）与其将来值之间的关系，初始投资获得的收益率（每期的利率）表示为 r ，其将来值（future value, FV）将于 N 年或 N 期后收到。

下面的例子说明了这个概念。假设你将 100 美元（ $PV = \$100$ ）投资到一个每年支付 5% 利息的银行账户中。在第 1 年年末，你将会获得 100 美元加上所得的利息， $0.05 \times \$100 = \5 ，一共是 105 美元。为将这个单期的例子用公式表示，我们定义以下这些符号：

- PV 表示投资的现值；
- FV_N 表示投资在 N 期后的将来值；
- r 表示每期的利率。

对于 $N = 1$ （即单期），现值 PV 的将来值表达式为

$$FV_1 = PV(1 + r) \tag{1-1}$$

对于上述这个例子，我们将 1 年后的将来值计算为 $FV_1 = \$100 \times (1.05) = \105 。

现在假设你决定将最初的 100 美元投资两年，且将每年所得利息存入账户（年复利）。在第 1 年年末（即第 2 年年初），你的账户将会拥有 105 美元，这 105 美元将会继续在你的银行账户中投资 1 年。于是，第 2 年年初的 105 美元（ $PV = \$105$ ）将会在第 2 年年末变为 $\$105(1.05) = \110.25 。注意，在第 2 年年末获得的 5.25 美元是第 2 年年初金额的 5%。

我们也可以这样来理解这个例子：第 2 年年初投资的金额是由最初的 100 美元与第 1 年获得的利息 5 美元组成的，而在第 2 年中，最初的 100 美元本金再一次获得利息，第 1 年获得的利息也同样获得利息。你可以看到最初的投资是如何增长的：

原始投资	\$100.00
第 1 年的利息（ $\$100 \times 0.05$ ）	5.00
基于原始投资的第 2 年利息（ $\$100 \times 0.05$ ）	5.00
基于原始投资利息的第 2 年利息（ $0.05 \times \$5.00$ 利息的利息）	0.25
总计	\$110.25

由 100 美元原始投资在每期获得的 5 美元利息被称为是单利（simple interest，利率乘以本

金)。本金 (principal) 是最初投资的资金额。在两年的期限中, 你共获得了 10 美元的单利。而在第 2 年年末你获得的额外的 0.25 美元, 是由你第 1 年获得的 5 美元利息再投资而产生的。

由利息再产生的利息让我们对复利 (compounding) 现象有了初步的认识。虽然初始投资所带来的利息是重要的, 但是在给定利率的情况下, 这笔利息在每一期中都是固定不变的。由利息再投资而获得的复利会产生一个强大得多的效果, 因为, 对于给定的利率, 它的规模每期都在增长, 复利的重要性随着利率量级的增长而增长。例如, 假设以 5% 的年利率进行复利计息, 那么今天投资的 100 美元, 将会在 100 年后变成 13 150 美元; 如果我们以 13% 的利率按年复利计息的话, 在 100 年后将会超过 2 000 万美元。

为了证实 2 000 万美元这个数字结果, 我们需要一个一般性的公式来处理任何期数的复利。下面的一般性的公式将初始投资的现值与其 N 期后的将来值联系起来:

$$FV_N = PV(1 + r)^N \tag{1-2}$$

式中, r 表示给定的每期利率; N 表示复利的期数。在上述银行账户投资的例子中, $FV_2 = \$100 \times (1 + 0.05)^2 = \110.25 。在以 13% 利率投资的例子中, $FV_{100} = \$100 \times (1.13)^{100} = \$20\,316\,287.42$ 。

使用将来值等式需要记住的最重要的一点是, 等式中的利率 r 和复利期数 N 必须相互匹配。这两个变量必须以相同的时间单位进行定义。例如, 如果 N 是以月作为单位的, 那么 r 就必须是 1 个月的利率, 而不是年利率。

一根时间轴可以帮助我们了解时间单位与每期的利率是否相互匹配。在时间轴上, 我们用时间指标 t 代表距离今天指定期数的一个时点。这样, 现值就是当前 (标注为 $t = 0$) 可用于投资的金额。我们现在可以将距离当前 N 期的时点标注为 $t = N$ 。图 1-1 所示的时间轴反映了这一关系。

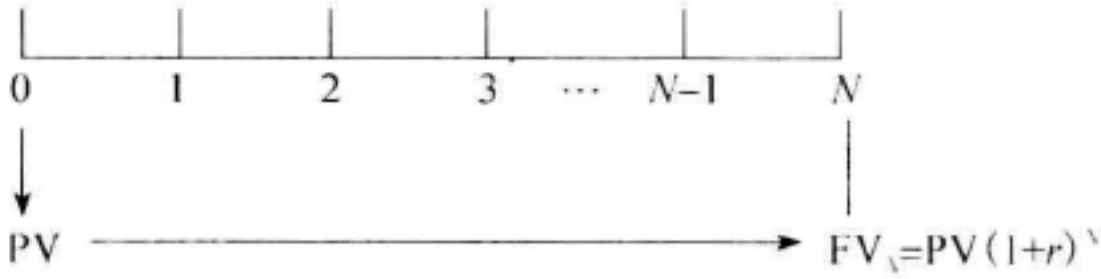


图 1-1 初始投资、现值 PV 及其将来值 FV 之间的关系

在图 1-1 上, 我们在 $t = 0$ 时刻标注了初始投资 PV。使用式 (1-2), 我们将现值乘以因子 $(1 + r)^N$ 使之向前移动到 $t = N$ 时刻。该因子称为将来值因子。我们把将来值在时间轴上表示为 FV, 并将其标注在 $t = N$ 时刻。假设将来值正好在 10 个期间后获得 ($N = 10$), 通过把现值 PV 乘以将来值因子 $(1 + r)^{10}$, 使现值 PV 和将来值 FV 处于时间轴的不同位置上。

现值和将来值在时间上不同, 由此可得到如下重要的结论:

- 我们只能对标注在同一时点上的货币金额进行加总。
- 给定利率, 将来值会随着期间数的增加而增加。
- 给定期数, 将来值会随着利率的增加而增加。

为了更好地理解这些概念, 可以考虑下面三个例子, 它们说明了应该如何应用将来值公式。

例 1-1 区间现金流按相同利率再投资的一次性投资的将来值

假如你幸运地成了州发行的税后 500 万美元彩票的中奖者。你决定将获得的奖金投资到当地金融机构的一个 5 年期的大额存单 (CD) 之中。该存单承诺每年支付 7% 的按年复利的利息。这家机构还允许你在这个存单的存续期中将利息以相同的利率再投资。如果你的资金始终按 7% 的年复利投资于这个账户, 期间从未取出, 那么 5 年后, 你将获得多少资金?

解：为了求解这个问题，我们使用式（1-2）来计算 500 万美元投资的将来值：

$$PV = \$5\,000\,000$$

$$r = 7\% = 0.07$$

$$N = 5$$

$$\begin{aligned} FV_N &= PV(1 + r)^N \\ &= \$5\,000\,000 \times (1.07)^5 \\ &= \$5\,000\,000 \times (1.402\,552) \\ &= \$7\,012\,758.65 \end{aligned}$$

在第 5 年年末，如果你的资金始终投资于这个 7% 的账户中而没有取出，你将拥有 7 012 758.65 美元。

在本例以及本章大部分的例子中，我们要注意，虽然将来值因子的计算值只保留到小数点后 6 位，但是所进行的计算实际上可能采用了更高的精度。例如，上面所显示的 1.402 552 是由 1.402 551 73 四舍五入而得（这一计算实际上是通过计算器或者电子表格采用至少小数点后 8 位数的精度完成的）。我们最终的结果是由计算器或者电子表格计算结果四舍五入而得到的。^①

例 1-2 没有区间现金流的一次性投资的将来值

一家机构提供给你如下的合约条款：对于一项 2 500 000 美元的投资，该机构承诺在 6 年后一次性以 8% 的年利率给予你支付。那么你预期在未来可以获得多少金额？

解：将如下数据代入式（1-2）中求解将来值：

$$PV = ¥2\,500\,000$$

$$r = 8\% = 0.08$$

$$N = 6$$

$$\begin{aligned} FV_N &= PV(1 + r)^N \\ &= ¥2\,500\,000 \times (1.08)^6 \\ &= ¥2\,500\,000 \times (1.586\,874) \\ &= ¥3\,967\,186 \end{aligned}$$

6 年后你预期获得 3 967 186 美元。

我们的第三个例子是一个更加复杂的将来值问题，它阐明了了解实际日历时间的重要性。

例 1-3 未来一次性支付款项的将来值问题

一位养老基金经理估计他公司的出资人在 5 年后将会出资 1 000 万美元，而该计划出资的资产收益率估计为每年 9%。这个养老基金的经理想要计算 15 年后该项投资的将来值，到那时该基金将把收益分配给退休投资者。这个将来值是多少呢？

解：通过确定初始投资 PV 的时点位于 $t = 5$ 时刻，我们可以将下面的数据代入式（1-2）来计算该项投资的将来值：

① 我们也可以利用利率因子表来求解货币时间价值的问题。使用利率因子表中的数据计算，其结果一般没有计算器或电子表格所计算的结果精确。因此，金融从业者更喜欢用计算器或电子表格进行计算。

$$\begin{aligned} PV &= \$10\,000\,000 \\ r &= 9\% = 0.09 \\ N &= 10 \\ FV_N &= PV(1+r)^N \\ &= \$10\,000\,000 \times (1.09)^{10} \\ &= \$10\,000\,000 \times (2.367\,364) \\ &= \$23\,673\,636.75 \end{aligned}$$

这个问题看上去很像前面的两个问题，但是它在一个重要的方面与之前的例子不同：它的投资时间。从今天（ $t=0$ ）的立场上来看，23 673 636.75 美元的将来金额是在未来 15 年后获得的。虽然将来值距离它的现值是 10 年，但是这 1 000 万美元的现值在未来的 5 年内是无法获得的。

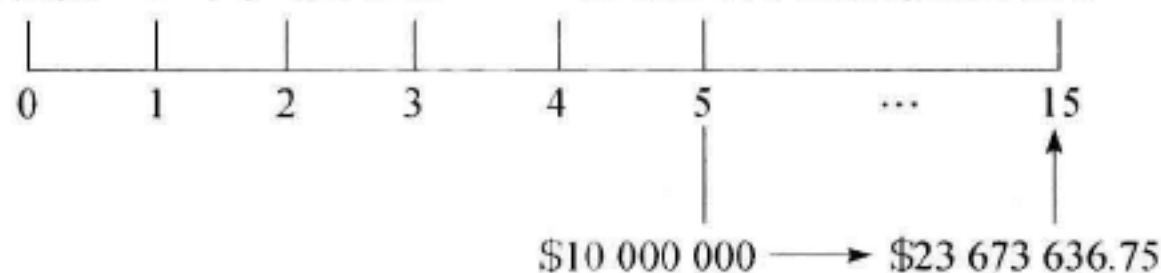


图 1-2 一次性投资的将来值，初始投资不在 $t=0$ 时刻

如图 1-2 所示，我们按照惯例将今天标示为 $t=0$ 时刻，而将之后的时间以 1 的间隔进行标注。这项 1 000 万美元的额外投资是在 5 年后收到的，所以它应该标示在 $t=5$ 时刻，出现在图 1-2 所示的位置上。于是，10 年后该项投资的将来值就应该标示在 $t=15$ 时刻；也就是在 $t=5$ 时刻收到 1 000 万美元投资的 10 年后。当我们处理更加复杂的问题时，像这种时间轴的标注方法是非常有用的，尤其是当问题涉及不止一笔现金流的时候。

在本章后面的一节中，我们将会讨论如何来计算 5 年之后获得的 1 000 万美元的投资在今天的价值问题。此处，我们暂且使用式 (1-2) 略作计算。假设上面的例 1-3 中的养老基金经理今天从公司出资人处收到 6 499 313.86 美元。那么在 5 年后这笔钱将值多少？在 15 年后又值多少呢？

$$\begin{aligned} PV &= \$6\,499\,313.86 \\ r &= 9\% = 0.09 \\ N &= 5 \\ FV_N &= PV(1+r)^N \\ &= \$6\,499\,313.86 \times (1.09)^5 \\ &= \$6\,499\,313.86 \times (1.538\,624) \\ &= \$10\,000\,000 \text{ 在 5 年标注的地方} \end{aligned}$$

并且

$$\begin{aligned} PV &= \$6\,499\,313.86 \\ r &= 9\% = 0.09 \\ N &= 15 \\ FV_N &= PV(1+r)^N \\ &= \$6\,499\,313.86 \times (1.09)^{15} \\ &= \$6\,499\,313.86 \times (3.642\,482) \\ &= \$23\,673\,636.74 \text{ 在 15 年标注的地方} \end{aligned}$$

这些结果表明今天 650 万美元左右的现值将在 5 年后变为 1 000 万美元, 在 15 年后变为 2 367 万美元。

1.3.1 复利的频数

在本节中, 我们考察在 1 年中不只付息一次的投资问题。例如, 许多银行提供在 1 年中复利 12 次的月度利率。在这样的安排下, 银行每个月会支付利息的利息。金融机构通常会以年利率报价, 而不是以周期性的月度利率报价, 这个年利率称为报价 (年) 利率 (stated annual interest rate or quoted interest rate)。我们将报价年利率记为 r_s 。例如, 你的银行可能会为一个特定的大额存单报出支付 8% 按月复利的利率。这个报价年利率等于月度利率乘以 12。在这个例子中, 月度利率为 $0.08/12 = 0.0067$ 或者是 0.67% ^①。这个利率是一个严格的报价惯例, 因为 $(1 + 0.0067)^{12} = 1.083$, 不是 1.08; 当复利次数比每年一次更加频繁时, $(1 + r_s)$ 项并不意味着是一个将来值因子。

对于每年不止复利一次的情况, 将来值公式可以表述为:

$$FV_N = PV \left(1 + \frac{r_s}{m} \right)^{mN} \quad (1-3)$$

式中, r_s 为报价年利率; m 为每年复利的次数; N 为年数。

请注意此处所使用的利率 r_s/m 与复利期数 mN 之间的匹配性。复利周期利率 r_s/m 是由报价年利率除以每年的复利期数得到的。复利期数 mN 是由一年中复利期数乘以年数得到的。复利周期利率 r_s/m 与复利期数 mN 必须匹配。

例 1-4 按季度复利的一次性投资的将来值

接着前面的大额存单的例子, 假设你的银行提供你一份两年到期的大额存单, 其报价年利率为 8%, 按季度复利。另外, 该投资的利息能够以相同的利率进行再投资。如果你决定投资 10 000 美元, 那么这份大额存单到期将值多少钱?

解: 利用式 (1-3) 计算将来值如下:

$$\begin{aligned} PV &= \$10\,000 \\ r_s &= 8\% = 0.08 \\ m &= 4 \\ r_s/m &= 0.08/4 = 0.02 \\ N &= 2 \\ mN &= 4 \times 2 = 8 \text{ 计息次数} \\ FV_N &= PV \left(1 + \frac{r_s}{m} \right)^{mN} \\ &= \$10\,000(1.02)^8 \\ &= \$10\,000(1.171\,659) \\ &= \$11\,716.59 \end{aligned}$$

到期时, 这份大额存单将值 11 716.59 美元。

① 当使用金融计算器时, 为了避免四舍五入带来的误差, 在 8 除以 12 后按 %i 键, 而不是简单地输入 0.67 及 %i, 由此我们得到 $(1 + 0.08/12)^{12} = 1.083\,000$ 。

计算将来值的式(1-3)与式(1-2)并没有本质区别。我们只要简单地记住所使用的利率是每个复利周期的利率,而所用的幂次是总的付息次数,或者说是复利期数。

例 1-5 按月度复利的一次性投资的将来值

一家澳大利亚的银行提供你一份利息6%、按月度复利支付的投资产品。你决定用100万澳元在该产品上投资1年。如果所有利息支付都能以6%的利率进行再投资,那么你的这项投资的将来值会是多少?

解:我们可以使用式(1-3)来求解这个1年期投资的将来值,具体计算如下:

$$\begin{aligned} PV &= \text{A\$}1\,000\,000 \\ r_s &= 6\% = 0.06 \\ m &= 12 \\ r_s/m &= 0.06/12 = 0.005\,0 \\ N &= 1 \\ mN &= 12(1) = 12 \text{ 计息期数} \\ FV_N &= PV \left(1 + \frac{r_s}{m} \right)^{mN} \\ &= \text{A\$}1\,000\,000(1.005)^{12} \\ &= \text{A\$}1\,000\,000(1.061\,678) \\ &= \text{A\$}1\,061\,677.81 \end{aligned}$$

如果你按6%年度复利计算,那么该投资的将来值仅为 $\text{A\$}1\,000\,000(1.06) = \text{A\$}1\,060\,000$,而不是按月度复利的1 061 677.81 澳元。

1.3.2 连续复利

在此之前关于复利期数的讨论给出了按离散时间间隔计算利息的离散复利的计算方法。如果每年复利期数变成无限多,那么这样的计息方式就被称为连续复利。如果我们想要对连续复利使用将来值公式,就需要求解式(1-3)中将来值因子在 $m \rightarrow \infty$ (每年复利期数无穷多)时的极限值。于是,N年以连续复利计息的一次性投资的将来值的数学表述为:

$$FV_N = PVe^{r_s N} \quad (1-4)$$

$e^{r_s N}$ 这项是指超越数 $e \approx 2.718\,281\,8$ 的 $r_s N$ 次方。大部分金融计算器都含有函数 e^x 。

例 1-6 按连续复利计息的一次性投资的将来值

假设一笔10 000美元的两年期投资将获得8%的连续复利。我们可以利用式(1-4)计算其将来值如下:

$$\begin{aligned} PV &= \$10\,000 \\ r_s &= 8\% = 0.08 \\ N &= 2 \end{aligned}$$

$$\begin{aligned}FV_N &= PVe^{r_sN} \\&= \$10\,000e^{0.08(2)} \\&= \$10\,000(1.173\,511) \\&= \$11\,735.11\end{aligned}$$

与例 1-4 中按季度计息的 11 716.59 美元相比，这里虽然使用了相同的利率，但是它是连续复利进行计息的，这一变化使得 10 000 美元的投资在两年内增长到了 11 735.11 美元。

表 1-1 列举了以 8% 报价年利率的一项 1 美元的初始投资分别按年度、半年度、季度、月度、日和连续复利计息，在 1 年年末产生不同数量美元的情况。（结果保留到小数点后 6 位）

表 1-1 复利频率对于将来值的影响

频率	r_s/m	mN	1 美元的将来值
年度	$8\%/1 = 8\%$	$1 \times 1 = 1$	$\$1.00(1.08) = \1.08
半年度	$8\%/2 = 4\%$	$2 \times 1 = 2$	$\$1.00(1.04)^2 = \$1.081\,600$
季度	$8\%/4 = 2\%$	$4 \times 1 = 4$	$\$1.00(1.02)^4 = \$1.082\,432$
月度	$8\%/12 = 0.666\,7\%$	$12 \times 1 = 12$	$\$1.00(1.006\,667)^{12} = \$1.083\,000$
日	$8\%/365 = 0.021\,9\%$	$365 \times 1 = 365$	$\$1.00(1.000\,219)^{365} = \$1.083\,278$
连续复利			$\$1.00e^{0.08(1)} = \$1.083\,287$

正如表 1-1 所示，所有 6 种情况都具有相同的 8% 的报价年利率，但是它们的期末美元数量却因为不同的复利频率而不同。按年度复利计息，其期末值为 1.08 美元。更加频繁的复利会产生更大的期末数值。按连续复利计息的期末美元数量在所有报价年利率为 8% 的投资中所获得的收益最大。

表 1-1 同样也显示了按 8.16% 年度复利计息的 1 美元投资将会在 1 年年末增长到与按 8% 半年度复利计息的 1 美元投资相同的将来值水平。这一结果使得我们了解到报价年利率与有效年利率（effective annual rate, EAR）之间的区别。^①对于按半年度复利的 8% 的报价年利率来说，其有效年利率为 8.16%。

1.3.3 报价利率和有效利率

因为报价年利率并不直接给出将来值，所以我们需要一个反映其有效年利率（EAR）的公式。对于按半年度复利计息的 8% 的年利率，我们将在计息期间（每半年）中以 4% 的利率计息。在 1 年的期间中，1 美元的投资将增长到 $\$1(1.04)^2 = \$1.081\,6$ ，如表 1-1 所示。该 1 美元的投资所获得的利息为 0.081 6 美元，它代表了 8.16% 的有效年利率。有效年利率可以按如下公式计算：

$$EAR = (1 + \text{期间利率})^m - 1 \tag{1-5}$$

期间利率是报价年利率除以 m ，其中 m 为 1 年中复利期间数。利用表 1-1，我们可以求得有效年利率为： $(1.04)^2 - 1 = 8.16\%$ 。

有效年利率的概念可以用到连续复利上。假设我们有一个 8% 的连续复利利率。我们可以用与上面相同的方式寻找适当的将来值因子，来求得其有效年利率。在这个例子中，一个 1 美元的

① 付息银行存款的有效年收益，在美国用的术语是年百分比收益（annual percentage yield, APY），而在英国用的是等价年利率（equivalent annual rate, EAR）。相反，年度百分比利率（annual percentage rate, APR）度量的是以年利率表示的借款成本。在美国，年百分比利率是通过期间利率乘以每年付息次数的数目得到的。因此，一些作者常使用年百分比利率作为报价年利率通常的同义词。然而，年百分比利率是一个具有法定含义的术语；它的计算遵从各国所各自规定的标准。因此，“报价年利率”是年利率（在一年中不进行复利计息的）首选的最一般的术语表述。

投资将会增长到 $\$1e^{0.08(1)} = \1.0833 。在这一年中所获得的利息代表有效年利率 8.33%，它要比半年度复利计息的有效年利率 8.16% 大，因为利息的复利更加频繁了。在连续复利的情况下，我们可以求解得到如下有效年利率：

$$\text{EAR} = e^{r_s} - 1 \quad (1-6)$$

我们可以逆向使用离散复利和连续复利的有效年利率公式，来得到与给定有效年利率相对应的期间利率。假设我们想要知道给定半年度复利计息的 8.16% 的有效年利率所对应的确切的期间利率是多少，那么我们可以利用式 (1-5) 来对期间利率进行求解：

$$0.0816 = (1 + \text{期间利率})^2 - 1$$

$$1.0816 = (1 + \text{期间利率})^2$$

$$(1.0816)^{1/2} - 1 = \text{期间利率}$$

$$1.04 - 1 = \text{期间利率}$$

$$4\% = \text{期间利率}$$

要计算与有效年利率 8.33% 相对应的连续复利利率（以连续复利计息的报价年利率），我们可以去寻找满足式 (1-6) 的利率：

$$0.0833 = e^{r_s} - 1$$

$$1.0833 = e^{r_s}$$

要解这个方程，我们可以在两边同时取自然对数（回忆一下， e^{r_s} 的自然对数是 $\ln e^{r_s} = r_s$ ）。因此， $\ln 1.0833 = r_s$ ，最终求得 $r_s = 8\%$ 。于是，我们了解到以连续复利计息的 8% 的报价年利率与 8.33% 的有效年利率是等价的。

1.4 现金流序列的将来值

在这一节中，我们考察现金流序列，包括等额的与不等额的两种类型。我们将先介绍一组在对分布于多个时间期间的现金流进行估值计算时常用的术语。

- 年金（annuity）是指一组有限的等额现金流序列。
- 普通年金（ordinary annuity）是指首笔现金流发生在一个期间段之后（标注为时点 $t=1$ ）的年金。
- 预付年金（annuity due）是指首笔现金流在当前即刻（标注为时点 $t=0$ ）发生的年金。
- 永续年金（perpetuity）是指无限期的年金，或者一组永不终结的等额现金流序列，其首笔现金流发生在一个期间段之后。

1.4.1 等额现金流序列——普通年金

考虑一个每年以 5% 利率支付的普通年金。假设我们先后有 5 笔等时间间隔（1 年）发生的 1 000 美元存款，其中首笔支付发生在 $t=1$ 时刻。我们的目标是求解这组普通年金在 $t=5$ 时刻（最后一笔存款发生后）的将来值。这里，考虑的时间间隔是 1 年，因此最后一笔支付发生在 5 年之后。如图 1-3 的时间轴上所示，我们可以利用式 (1-2) $[FV_N = PV(1+r)^N]$ 来得到每笔 1 000 美元存款在时点 $t=5$ 上的将来值。图 1-3 中的箭头都是从支付时点指向 $t=5$ 时点。例如，首笔在 $t=1$ 时点发生的 1 000 美元存款将在 4 个期间上复利。使用式 (1-2)，我们可以计算得到其在 $t=5$ 时点的将来值为 $\$1\,000(1.05)^4 = \$1\,215.51$ 。用类似的方法，我们可以计算得到所有其他支付的

将来值。(注意,我们所要计算的是在 $t=5$ 时点上的将来值,因此最后一笔支付不会获得任何利息。)现在我们得到了所有支付在 $t=5$ 时点上的将来值,然后将这些将来值加总,从而得到该年金的将来值,其值为 5 525.63 美元。



图 1-3 5 年期普通年金的将来值

如果我们将年金支付额记为 A , 支付期数记为 N , 每期利率记为 r , 那么我们就可以得到一般的年金计算公式。我们可以定义将来值为:

$$FV_N = A[(1+r)^{N-1} + (1+r)^{N-2} + (1+r)^{N-3} + \cdots + (1+r)^1 + (1+r)^0]$$

该公式可以简化为:

$$FV_N = A \left[\frac{(1+r)^N - 1}{r} \right] \quad (1-7)$$

在括号中的项是将来值年金因子。这个因子给出了每期支付 1 美元的普通年金的将来值。将年金支付额乘以将来值年金因子就可以得到该组普通年金的将来值。对于图 1-3 中的普通年金来说,我们可以从式 (1-7) 中求得将来值年金因子为

$$\left[\frac{(1.05)^5 - 1}{0.05} \right] = 5.525631$$

当年金每期支付额 $A = \$1\,000$, 该年金的将来值为 $\$1\,000(5.525631) = \$5\,525.63$, 与我们之前的计算结果是一致的。

下面的例 1-7 说明了如何使用式 (1-7) 来计算普通年金的将来值。

例 1-7 年金的将来值

假设你公司的固定供款退休金计划允许你每年投资 20 000 欧元。你计划在未来的 30 年, 每年投资 20 000 欧元于股票指数基金之中。从历史上来看, 该基金平均每年将带来 9% 的收益。假设你每年确实能够获得 9% 的收益, 那么在你最后一次付款后可获得多少财富以备退休之需呢?

解: 利用式 (1-7) 来求解将来值:

$$A = \text{€}20\,000$$

$$r = 9\% = 0.09$$

$$N = 30$$

$$\text{将来值年金因子} = \frac{(1+r)^N - 1}{r} = \frac{(1.09)^{30} - 1}{0.09} = 136.307539$$

$$FV_N = \text{€}20\,000(136.307539)$$

$$= \text{€}2\,726\,150.77$$

假设基金每年能够持续地平均获益 9%, 那么你将在退休时拥有 2 726 150.77 欧元的财富。

1.4.2 不等额现金流序列

在许多情况下，现金流序列并不是等额的，这就使我们无法简单地利用将来值年金因子来进行计算。例如，某位个人投资者可能有一个储蓄计划，该计划会根据年度中月份的不同进行不等额的现金支付或者会在计划的度假期间减少储蓄额。我们总是可以通过每次计算其中一笔现金流的将来值来求得该不等额现金流序列的将来值。假设你有如表 1-2 所示的 5 笔现金流，所标注的时点是相对于当前 ($t=0$) 时点的。

表 1-2 一组不等额现金流和其在 5% 利率水平下的将来值

时点	现金流	在第 5 年的将来值
$t=1$	\$1 000	$\$1\,000\ (1.05)^4 = \$1\,215.51$
$t=2$	\$2 000	$\$2\,000\ (1.05)^3 = \$2\,315.25$
$t=3$	\$4 000	$\$4\,000\ (1.05)^2 = \$4\,410.00$
$t=4$	\$5 000	$\$5\,000\ (1.05)^1 = \$5\,250.00$
$t=5$	\$6 000	$\$6\,000\ (1.05)^0 = \$6\,000.00$
		总和 = \$19 190.76

所有的支付如表 1-2 所示是不同的。因此，求解在 $t=5$ 时点将来值的最直接的方法就是计算每笔支付在 $t=5$ 时刻的将来值，然后将所有单个的将来值进行加总。如表中第三列所示，在第 5 年年末总的将来值为 19 190.76 美元。在本章的后面部分，你将会学到当现金流接近等额时的快捷计算方法；这些快捷的计算方法会使你将年金计算与单期现金流计算相联系。

1.5 单笔现金流的现值

1.5.1 求解单笔现金流的现值

正如将来值因子把今天的现值与明天的将来值相联系，现值因子使我们把将来值折现到现值。例如，如果已知 5% 的利率在 1 年后会产生 105 美元的收益，那么我们在当前要投资多少资金才能于 1 年后以 5% 利率获得 105 美元？这个答案是 100 美元；因此，100 美元就是 1 年后获得的 105 美元以 5% 利率进行折现所获得的现值。

给定 N 期后所获得的未来现金流和单个期间的利率 r ，我们可以利用将来值公式直接求解其现值如下：

$$\begin{aligned} FV_N &= PV(1+r)^N \\ PV &= FV_N \left[\frac{1}{(1+r)^N} \right] \\ PV &= FV_N(1+r)^{-N} \end{aligned} \tag{1-8}$$

我们可以看到，式 (1-8) 中的现值因子 $1/(1+r)^N$ 就是将来值因子 $[(1+r)^N]$ 的倒数。

例 1-8 一次性投资的现值

一家保险公司发行了一个担保投资合同 (Guaranteed Investment Contract, GIC)，该合同承诺在 6 年后以 8% 的回报率支付 100 000 美元。那么，投保者现在必须投资多少资金才能以 8% 的利率在 6 年后获得该承诺的支付？

解：我们利用下列数据和式（1-8）来求解其现值：

$$FV_N = \$100\,000$$

$$r = 8\% = 0.08$$

$$N = 6$$

$$PV = FV_N(1 + r)^{-N}$$

$$= \$100\,000 \left[\frac{1}{(0.08)^6} \right]$$

$$= \$100\,000(0.630\,169\,6)$$

$$= \$63\,016.96$$

我们可以认为今天投资的 63 016.96 美元，在 8% 的利率水平下，等价于在 6 年后收到的 100 000 美元。当对于货币时间价值进行考虑时，未来的 100 000 美元经过贴现将与今天的 63 016.96 美元等价。如图 1-4 中的时间轴所示，100 000 美元被折现到 6 个完整期间之前的 0 时点。

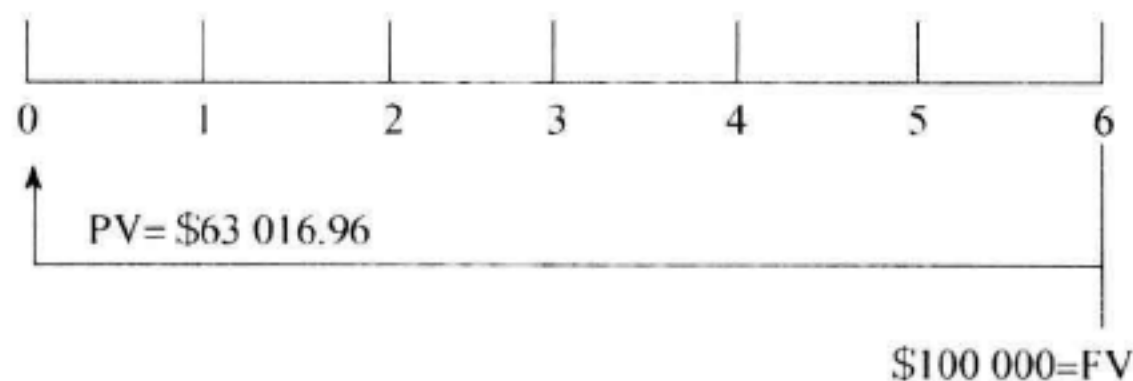


图 1-4 在时点 $t=6$ 收到的一次性支付的现值

例 1-9 更远未来的一次性投资的投影现值^①

假设你拥有一份流动性的金融资产，它将在 10 年之后支付给你 100 000 美元。由于你的女儿计划在 4 年后上大学，因此，你想知道该资产在那个时候的现值是多少。给定 8% 的贴现率，该份资产在 4 年后的价值是多少？

解：这份资产的价值是其承诺支付额的现值。在 $t=4$ 时点上，其现金支付将在 6 年后收到。有了这些信息，你可以利用式（1-8）求解 4 年后的资产价值：

$$FV_N = \$100\,000$$

$$r = 8\% = 0.08$$

$$N = 6$$

$$PV = FV_N(1 + r)^{-N}$$

$$= \$100\,000 \left[\frac{1}{(1.08)^6} \right]$$

$$= \$100\,000(0.630\,169\,6)$$

$$= \$63\,016.96$$

① 此处投影现值（projected present value）指的是将 10 年后的 100 000 美元折算为 4 年后的价值，而非折算为当前的价值。——译者注

图 1-5 中的时间轴上标出了在 $t=10$ 时刻收到的未来支付 100 000 美元。时间轴上同样也标出了其在 $t=4$ 时刻和 $t=0$ 时刻的价值。相对于 $t=10$ 时刻的支付，在 $t=4$ 时刻的金额是一个投影现值，而 $t=0$ 时刻的值是现值（即今天的价值）。

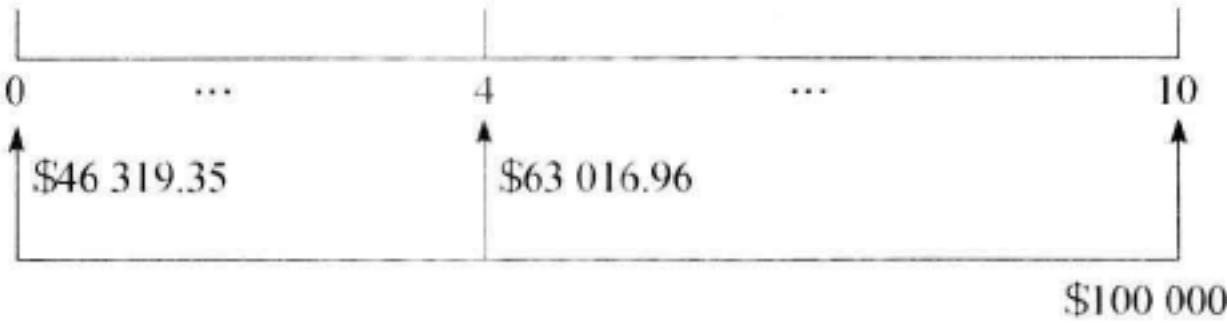


图 1-5 现值与将来值之间的关系

现值问题需要利用现值因子 $(1+r)^{-N}$ 对现金流进行估值计算。现值与贴现率和期数之间的关系表现为如下几个方面：

- 对于一个给定的贴现率，未来金额收到的时间越晚，该金额的现值就越小。
- 给定时间长度不变，贴现率越大，将来金额的现值就越小。

1.5.2 复利的频数

回忆一下，利息可以按半年度、季度、月度或者甚至按日支付。要处理在 1 年中不止支付一次利息的现值问题，我们需要调整现值公式 (1-8)。 r_s 是报价利率，等于期间利率乘以 1 年中复利次数。一般地，当 1 年中复利次数大于一次时，我们可以将现值公式表示为：

$$PV = FV_N \left(1 + \frac{r_s}{m} \right)^{-mN} \tag{1-9}$$

式中， m 表示每年的复利次数； r_s 表示报价年利率； N 表示年数。

式 (1-9) 与式 (1-8) 十分相似。正如我们已经注意到的，现值和将来值互为倒数。改变复利频率并不改变这一关系。两者唯一的区别在于期间利率和相应的复利次数的使用上。

下面举例说明式 (1-9) 的用法。

例 1-10 按月度复利计息的一次性支付的现值

加拿大一家养老基金的基金经理了解到该基金在 10 年后必须一次性支付 500 万加元。于是，她想要在今天投资一个担保投资合同，以使该基金到那时能增长到所需的金额。当前担保投资合同的利率为每年 6%，按月度复利计息。那么，该基金经理在今天应该投资多少资金于该担保投资合同？

解：利用式 (1-9) 来求解所要求的现值：

$$\begin{aligned} FV_N &= \text{C\$}5\,000\,000 \\ r_s &= 6\% = 0.06 \\ m &= 12 \\ r_s/m &= 0.06/12 = 0.005 \\ N &= 10 \\ mN &= 12 \times 10 = 120 \end{aligned}$$

$$\begin{aligned} PV &= FV_N \left(1 + \frac{r_s}{m}\right)^{-mN} \\ &= \text{C\$}5\,000\,000(1.005)^{-120} \\ &= \text{C\$}5\,000\,000(0.549\,633) \\ &= \text{C\$}2\,748\,163.67 \end{aligned}$$

在运用式(1-9)的过程中,我们使用了期间利率(在这个例子中是月利率)和适当的按月度复利计息的期数(在这个例子中,是按月度复利计息10年,即120个期间)。

1.6 现金流序列的现值

投资管理中的应用问题所涉及的资产常常提供一系列分布在不同时间上的现金流。现金流可能是差别极大的,也可能是相对等额的,或者是完全等额的。它们可能发生在较短的期间,也可能发生在较长的期间,甚至可能是延伸至无限期的。在这一节中,我们讨论如何求解一系列现金流的现值。

1.6.1 等额现金流序列的现值

首先,我们先从普通年金开始谈起。回忆一下,普通年金具有等额的年金支付,而且首笔支付发生在未来的第一期末。整个年金共有 N 期支付,首期发生在 $t=1$ 时刻,最后一期发生在 $t=N$ 时刻。我们可以将普通年金的现值表述为每个单笔年金支付的现值的总和,具体表示如下:

$$PV = \frac{A}{(1+r)} + \frac{A}{(1+r)^2} + \frac{A}{(1+r)^3} + \cdots + \frac{A}{(1+r)^{N-1}} + \frac{A}{(1+r)^N} \quad (1-10)$$

式中, A 表示(每期)年金支付额; r 表示与年金支付频率相关的每期利率(如按年、季或者月); N 表示年金支付期数。

因为年金支付额(A)在这个公式中是一个常数,它可以作为一个公因子提出,作为一个公共项。因此,利息因子的总和有一个简洁的表述:

$$PV = A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \quad (1-11)$$

与计算普通年金将来值的方式类似,我们可以将年金支付额乘以现值年金因子[式(1-11)括号中的项]来获得现值。

例 1-11 普通年金的现值

假设你正考虑购买一项承诺在未来5年内每年支付1 000欧元的金融资产,其首笔支付是在一年后提供,所要求的回报率为每年12%。那么你应该支付多少来购买这项资产?

解:我们利用式(1-11)所给出的普通年金现值的计算公式和如下的数据来求解该金融资产的价值:

$$\begin{aligned} A &= \text{€}1\,000 \\ r &= 12\% = 0.12 \\ N &= 5 \end{aligned}$$

$$\begin{aligned} PV &= A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \\ &= €1\,000 \left[\frac{1 - \frac{1}{(1.12)^5}}{0.12} \right] \\ &= €1\,000(3.604\,776) \\ &= €3\,604.78 \end{aligned}$$

在未来5年内每年支付1 000欧元的一系列现金流，以12%进行贴现的当前价值为3 604.78欧元。

现实中发生现金流的实际时间不同，我们考虑一种特定的等额年金：预付年金。

预付年金的首笔支付发生在今天 ($t=0$)。总体上，预付年金会有 N 次支付。图 1-6 在时间轴上描述了一个只有 4 次支付，每次支付额为 100 美元的预付年金。

如图 1-6 所示，我们可以将 4 期的预付年金视为两个部分的和：今天一次性支付的 100 美元和一个每期支付 100 美元共 3 期的普通年金。在 12% 的贴现率水平下，这个预付年金例子中的 4 个 100 美元的现金流将价值 340.18 美元。[⊖]



图 1-6 每期 100 美元的预付年金

将未来现金流序列的价值表达成当前美元的方法使得我们能够方便地将不同的年金进行比较。下面的例子就说明了这一方法。

例 1-12 即刻支付的现金流加上普通年金的现值

假如今天你要退休了，那么你必须选择获取你的退休金的方式，或者以一次性方式领取，或者以年金的方式领取。你所在公司的退休金管理人员提供你两个可选方案：一份即刻支付 200 万美元的支票；或者是一份每年支付 20 万美元、共支付 20 年的年金，该年金的首笔支付在今天即刻支付。目前，银行的利率为每年 7%，按年度复利。那么哪一个可选方案会有更大的现值？（忽略任何对于这两个选择的税收影响）

解：要比较这两个方案，我们就要求解每个方案在 $t=0$ 时刻的现值，然后选择其中具有较大现值的方案。第一个方案的现值是 200 万美元，因为它已经用今天的美元数量表示了。第二个方案是一个预付年金。因为首笔支付发生在 $t=0$ 时点，所以你可以将该年金收益分成两个部分：一个今天即刻支付（在 $t=0$ 时刻支付）的 200 000 美元以及一个每年支付 200 000 美元、共支付 19 期的普通年金。要计算该方案的价值，你需要用式 (1-11) 求出普通年金的现值，然后再加上 200 000 美元。

$$\begin{aligned} A &= \$200\,000 \\ N &= 19 \\ r &= 7\% = 0.07 \end{aligned}$$

⊖ 还有另外一个计算预付年金现值的方法。由于预付年金相对于普通年金，其每笔支付都要少一个单位期间，因此，我们可以对式 (1-11) 进行修改，即让等式右边乘以 $(1+r)$ 来对于预付年金进行处理。于是就有，
 $PV(\text{预付年金}) = A \{ [1 - (1+r)^{-N}] / r \} / (1+r)$ 。

$$\begin{aligned} PV &= A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \\ &= \$200\,000 \left[\frac{1 - \frac{1}{(1.07)^{19}}}{0.07} \right] \\ &= \$200\,000 (10.335\,595) \\ &= \$2\,067\,119.05 \end{aligned}$$

上述 19 期支付的 200 000 美元的现值为 2 067 119.05 美元，加上初始支付的 200 000 美元，我们可以得到该年金方案的现值合计为 2 267 119.05 美元。显然该年金的现值大于一次性支付 200 万美元的另一可选方案。

我们再来看另一个关于现值和将来值等价性的例子。

例 1-13 普通年金的投影现值

一位德国养老金基金经理预计每年要支付退休者 100 万欧元的退休金。退休金将在 10 年后 ($t=10$) 开始支付。一旦退休金开始支付，它将一直持续到 $t=39$ 时刻，共支付 30 期。那么，如果对于该养老金计划适当的年贴现率为 5%，按年度复利，则该养老金计划的现值是多少？

解：这个问题涉及首笔支付在 $t=10$ 时刻的年金。如果站在 $t=9$ 时刻的角度来看，我们会看到一个 30 期支付的普通年金。我们可以用式 (1-11) 来计算该年金的现值，然后在时间轴上来观察它。

$$\begin{aligned} A &= €1\,000\,000 \\ r &= 5\% = 0.05 \\ N &= 30 \\ PV &= A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \\ &= €1\,000\,000 \left[\frac{1 - \frac{1}{(1.05)^{30}}}{0.05} \right] \\ &= €1\,000\,000 (15.372\,451) \\ &= €15\,372\,451.03 \end{aligned}$$

在时间轴上，我们看到 1 000 000 欧元的养老金支付从 $t=10$ 时刻一直延续到 $t=39$ 时刻。大括号和箭头表示求解年金的过程，即将年金贴现到 $t=9$ 时点上。该养老金计划在 $t=9$ 时刻的现值为 15 372 451.03 欧元。现在的问题是求解在今天的现值（在 $t=0$ 时刻）。

现在我们可以依赖现值和将来值的等价关系来进行计算。如图 1-7 所示，我们可以将 $t=9$ 时点上的现值视为 $t=0$ 时点上的将来值。我们可以用下面的方式计算 $t=9$ 时点上金额的现值：

$$\begin{aligned}FV_N &= \text{€ } 15\,372\,451.03 \text{ (} t = 9 \text{ 时刻的现值)} \\N &= 9 \\r &= 5\% = 0.05 \\PV &= FV_N(1+r)^{-N} \\&= \text{€ } 15\,372\,451.03(1.05)^{-9} \\&= \text{€ } 15\,372\,451.03(0.644\,609) \\&= \text{€ } 9\,909\,219.00\end{aligned}$$

故该养老金计划的现值为 9 909 219.00 欧元。

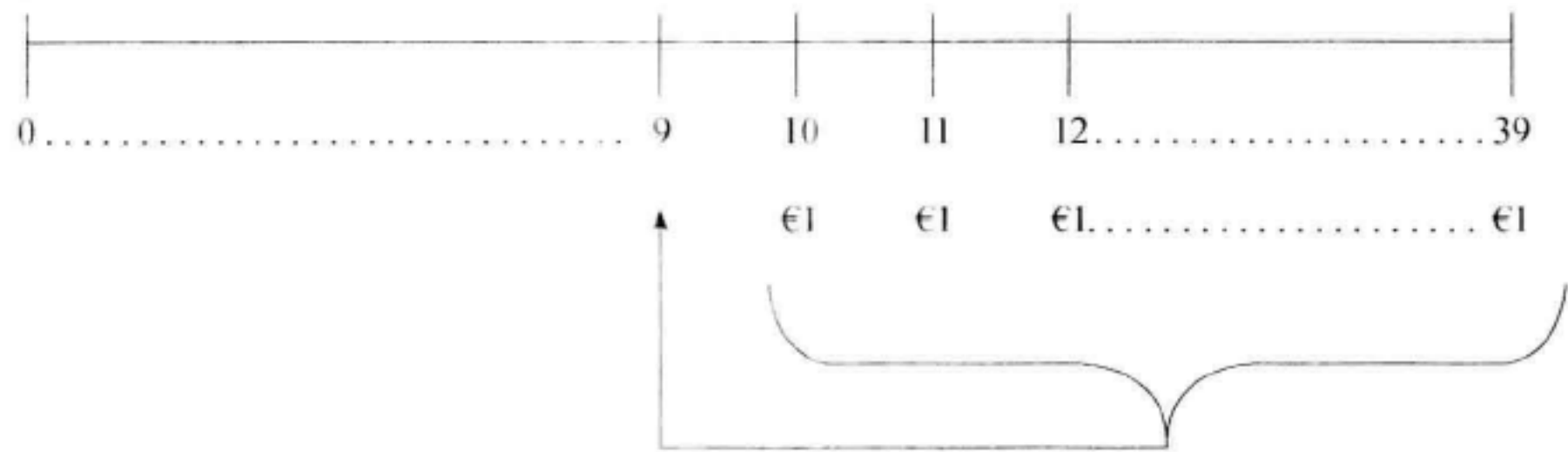


图 1-7 首笔支付在 $t = 10$ 时刻的普通年金的现值（以 100 万美元计）

例 1-13 演示了本章所强调的关于现金流计算的三个步骤：

- 求解任一现金流序列的现值或者将来值；
- 识别出现值与适当贴现后的将来值之间的等价性；
- 在货币时间价值计算中注意现金流发生的实际时间。

1.6.2 无限期等额现金流序列的现值——永续年金

考虑延伸至无限期的普通年金的情况。这种普通年金被称为永续年金。我们可以修改式 (1-10)，使其能够反映无限期的现金流序列，从而得到永续年金现值的计算公式：

$$PV = A \sum_{t=1}^{\infty} \left[\frac{1}{(1+r)^t} \right] \tag{1-12}$$

只要利率始终为正，该现值因子之和总是收敛的，其值为：

$$PV = \frac{A}{r} \tag{1-13}$$

要了解上式是如何得到的，我们可以回顾一下式 (1-11) 中关于普通年金现值的表达。当 N （年金支付的期数）趋向无穷， $1/(1+r)^N$ 这一项趋于 0，于是式 (1-11) 就简化为式 (1-13)。这个公式将在对股票红利流估值中再次出现，因为股票是没有事先确定的期限的。（一只支付固定红利的股票与永续年金是类似的。）一个每年支付 10 美元、首笔支付在 1 年后的永续年金，在 20% 的要求回报率条件下，其现值为 $\$10/0.2 = \50 。

式 (1-13) 只有在永续年金每期等额支付的情况下才是有效的。在我们上面的例子中，首笔支付发生在 $t = 1$ 时点上；因此，我们所计算的现值是在 $t = 0$ 时点上的。

其他资产同样可近似满足永续年金的假设。某些政府债券以及优先股就是具有无限期等额支付金融资产的典型例子。

例 1-14 永续年金的现值

英国政府曾经发行过一种叫做统一公债的证券,该证券承诺无限期地支付等额现金流。如果统一公债每年支付 100 英镑,那么,如果要求的回报率为 5%,则该证券在今天的价值为多少?

解:要回答该问题,我们可以利用式(1-13)和下列数据:

$$\begin{aligned}A &= £ 100 \\r &= 5\% = 0.05 \\PV &= A/r \\&= £ 100/0.05 \\&= £ 2\,000\end{aligned}$$

故该债券将价值 2 000 英镑。

1.6.3 始点不在零时刻的现金流序列的现值

在实际的投资实务中,分析师经常需要求解始点不在 $t=0$ 时刻的现值。通过改变现值公式中时点的下标并评估从第 2 年开始支付的每年 100 美元的永续年金的现值,我们可以得到 $PV_1 = \$100/0.05 = \$2\,000$ (以 5% 贴现率折现)。进一步,我们可以计算今天的现值为 $PV_0 = \$2\,000/1.05 = \$1\,904.76$ 。

考虑一个类似的情况,现金流序列从第 4 年年末开始,每年发生 6 美元的现金支付,直到第 10 年年末为止。从第 3 年年末的角度来看,我们所面对的是一个典型的 7 年期的普通年金。我们可以站在第 3 年年末的角度,求得该年金的现值,并将该现值贴现到当前的时刻。在 5% 利率水平下,该始于第 4 年年末、每年支付 6 美元的现金流序列在第 3 年年末将价值 34.72 美元 ($t=3$),而在今天将价值 29.99 美元 ($t=0$)。

下面的例子将说明一个重要的概念,即一个开始于未来某时刻的年金或永续年金可以表示成在其首笔支付时点的前一个时点上的现值。然后,该现值可以再被贴现到今天的现值。

例 1-15 投影永续年金的现值

考虑一个等额永续年金,其每年支付 100 英镑,首笔支付发生于 $t=5$ 时刻。给定 5% 的贴现率,该永续年金在今天 ($t=0$) 的现值为多少?

解:首先,我们求解永续年金在 $t=4$ 时刻的现值,然后将其贴现到 $t=0$ 时刻。(回忆一下,一个永续年金或者普通年金的首笔支付是发生在一期之后,所以我们在现值计算中选择时点为 $t=4$ 。)

1. 求解在 $t=4$ 时刻的永续年金的现值:

$$\begin{aligned}A &= £ 100 \\r &= 5\% = 0.05 \\PV &= A/r \\&= £ 100/0.05 \\&= £ 2\,000\end{aligned}$$

2. 求解在 $t=4$ 时刻一笔未来金额的现值。从 $t=0$ 时点的角度来看, 之前计算的现值 £2 000 可以被看做是一个将来值。现在我们需要求解该一次性支付的现值:

$$FV_N = £2\,000 \text{ (在 } t=4 \text{ 时刻的现值)}$$

$$r = 5\% = 0.05$$

$$N = 4$$

$$\begin{aligned} PV &= FV_N(1+r)^{-N} \\ &= £2\,000(1.05)^{-4} \\ &= £2\,000(0.822\,702) \\ &= £1\,645.40 \end{aligned}$$

该永续年金在今天的现值为 1 645.40 英镑。

正如之前所讨论的, 一个年金是给定期数的一系列固定金额的支付。假设我们已经拥有了一个永续年金; 同时, 我们再发行一个由我们承担支付义务的永续年金, 该永续年金的支付额与我们所持有的永续年金的支付额相同, 但其支付是从 $t=5$ 时刻开始, 并一直持续到永远的。那么, 第二个永续年金的支付完全抵消了 $t=5$ 时刻以及在此之后所有的由我们所持有的永续年金所给予的支付。于是, 我们只剩下了 $t=1, 2, 3$ 和 4 时点上的非零等额现金流。这一结果与 4 期年金的定义完全一致。因此, 我们可以通过将两个支付相等但是具有不同起始点的永续年金作差来构造一个普通年金。下面的例子就说明了这一结果。

例 1-16 普通年金的现值等于当前开始支付的永续年金与未来某时点开始支付的投影永续年金之差的现值

给定 5% 的贴现率, 求解一个 4 年期普通年金的现值, 该年金每年支付 100 英镑, 起始支付在第 1 年年末。该年金可以视为下面两个等额永续年金之差:

永续年金 1 从第 1 年年末开始, 每年支付 100 英镑 (首笔支付在 $t=1$ 时刻)

永续年金 2 从第 5 年年末开始, 每年支付 100 英镑 (首笔支付在 $t=5$ 时刻)

解: 如果将第一个永续年金减去第二个永续年金, 我们将得到每期支付 100 英镑的 4 年期普通年金 (支付发生在 $t=1, 2, 3, 4$ 时刻)。通过将第一个永续年金的现值减去第二个永续年金的现值, 我们可以得到 4 年期普通年金的现值:

$$1. PV_0(\text{永续年金 } 1) = £100/0.05 = £2\,000$$

$$2. PV_4(\text{永续年金 } 2) = £100/0.05 = £2\,000$$

$$3. PV_0(\text{永续年金 } 2) = £2\,000/(1.05)^4 = £1\,645.40$$

$$\begin{aligned} 4. PV_0(\text{年金}) &= PV_0(\text{永续年金 } 1) - PV_0(\text{永续年金 } 2) \\ &= £2\,000 - £1\,645.40 \\ &= £354.60 \end{aligned}$$

故 4 年期普通年金的现值等于 $£2\,000 - £1\,645.40 = £354.60$ 。

1.6.4 不等额现金流序列的现值

当我们遇到非等额的现金流序列时, 我们必须首先求得每笔现金流的现值, 然后将各个现值

进行加总。对于一个具有许多笔现金流的序列，我们通常会使用电子表格来进行现值计算。表 1-3 列出了一个现金流序列，其中的第一列是现金流发生的时间点，第二列是现金流发生的金额，每笔现金流的现值位于第三列。表 1-3 最后一行反映的是这 5 笔现金流现值的总和。

表 1-3 一组不等额现金流及其以 5% 折现的现值

时间期间	现金流	在第 0 年的现值
1	\$1 000	$\$1\,000(1.05)^{-1} = \952.38
2	\$2 000	$\$2\,000(1.05)^{-2} = \$1\,814.06$
3	\$4 000	$\$4\,000(1.05)^{-3} = \$3\,455.35$
4	\$5 000	$\$5\,000(1.05)^{-4} = \$4\,113.51$
5	\$6 000	$\$6\,000(1.05)^{-5} = \$4\,701.16$
		总和 = \$15 036.46

我们可以通过利用单笔支付的将来值公式来逐一计算这些现金流的将来值。由于我们已经知道该序列的现值，所以我们可以方便地应用时间价值等价关系。表 1-2 中的这组现金流的将来值为 19 190.76 美元，这等价于当前单笔 15 036.46 美元的支付复利计息到 $t = 5$ 时刻的价值：

$$\begin{aligned} PV &= \$15\,036.46 \\ N &= 5 \\ r &= 5\% = 0.05 \\ FV_N &= PV(1 + r)^N \\ &= \$15\,036.46(1.05)^5 \\ &= \$15\,036.46(1.276\,282) \\ &= \$19\,190.76 \end{aligned}$$

1.7 求解利率、期数或年金支付额

在之前的例子中，相关的信息是已知的。例如，所有的问题都给定利率 r 、期数 N 、年金（每期）支付额 A 以及现值 PV 或者将来值 FV 。然而在现实世界的应用中，虽然现值和将来值可能是已知的，但是你可能不得不去求解利率、期数或者年金支付额。在下面的小节中，我们将讨论这类问题。

1.7.1 求解利率和增长率

假设我们已知一份 100 欧元的银行存款在 1 年后将会获得 111 欧元的支付。在这个信息下，我们可以利用式 (1-2)， $FV_N = PV(1 + r)^N$ ，式中 $N = 1$ ，推断出将现值 100 欧元从将来值 111 欧元中分解出来的利率。在 PV 、 FV 和 N 已知的情况下，我们可以直接求解 r ：

$$\begin{aligned} 1 + r &= FV/PV \\ 1 + r &= €111/€100 = 1.11 \\ r &= 0.11 \text{ 或者 } 11\% \end{aligned}$$

故使得 $t = 0$ 时刻的 100 欧元与 $t = 1$ 时刻的 111 欧元等价的利率为 11%。因此，我们可以认为 100 欧元在 11% 的增长率下会增长到 111 欧元。

正如这个例子所反映的，利率可以被认为是增长率。对于具体的应用，我们会决定采用术语“利率”或“增长率”。利用式 (1-2) 来求解 r ，并将利率 r 用增长率 g 来代替，于是就有了用来确定增长率的如下表述：

$$g = (FV_N/PV)^{1/N} - 1 \tag{1-14}$$

下面的两个例子将会用到上述增长率的概念。

例 1-17 增长率的计算 (1)

布兰兹 (Brands) 有限公司在 1998 年的净销售额为 8 436 000 000 美元。而在 2002 年的净销售额为 8 445 000 000 美元，只稍稍高于 1998 年。在从 1998 年年末至 2002 年年末的 4 年间，布兰兹有限公司净销售额的增长率为多少？

解：要解决这个问题，我们可以利用式 (1-14)， $g = (FV_N/PV)^{1/N} - 1$ 。我们记 1998 年的净销售额为 PV，2002 年的净销售额为 PV_4 。于是，我们可以求解增长率如下：

$$\begin{aligned} g &= \sqrt[4]{\$8\,445/\$8\,436} - 1 \\ &= \sqrt[4]{1.001\,067} - 1 \\ &= 1.000\,267 - 1 \\ &= 0.000\,267 \end{aligned}$$

所计算得出的增长率约为每年 0.03%，仅仅比零高出一点点。这和我们最初对于布兰兹有限公司净销售额在 1998 ~ 2002 年间基本平稳的直观印象是一致的。

例 1-18 增长率的计算 (2)

在例 1-17 中，我们发现布兰兹有限公司在 1998 ~ 2002 年期间的净销售额复利增长率接近于零。作为一个零售商，布兰兹有限公司的销售不仅依赖于其门店数量（或销售门店面积（平方英尺或平方米）），还依赖于其每个门店的销售额（或平均每单位销售门店面积（平方英尺或平方米）的销售额）。事实上，布兰兹有限公司在 1998 ~ 2002 年期间减少了其门店数量。在 1998 年，其门店数量为 5 382 家，而在 2002 年为 4 036 家。在这个例子中，我们会谈及一个正的复利的减少率或者说是一个负的复利增长率。那么，运营门店数量的增长率到底是多少呢？

解：使用式 (1-14)，我们可以求得：

$$\begin{aligned} g &= \sqrt[4]{4\,036/5\,382} - 1 \\ &= \sqrt[4]{0.749\,907} - 1 \\ &= 0.930\,576 - 1 \\ &= -0.069\,424 \end{aligned}$$

故在 1998 ~ 2002 年期间运营门店的增长率约为 -6.9%。需要注意的是，我们同样可以将 -6.9% 看做是复利年度增长率，因为该数字反映了门店数量从 1998 ~ 2002 年的复利增长情况。表 1-4 列出了 1998 ~ 2002 年期间布兰兹有限公司运营门店数量的变化情况。

表 1-4 布兰兹有限公司运营门店的数量 (1998 ~ 2002 年)

年份	门店数量	$(1 + g)_t$	t
1998	5 382		0
1999	5 023	$5\,023/5\,382 = 0.933\,296$	1
2000	5 129	$5\,129/5\,023 = 1.021\,103$	2
2001	4 614	$4\,614/5\,129 = 0.899\,591$	3
2002	4 036	$4\,036/4\,614 = 0.874\,729$	4

资料来源：www.limited.com.

表 1-4 给出了各年度的 $(1 + \text{门店数量年增长率})$ 的数值。将历年的 $(1 + \text{门店数量年增长率})$ 的数值相乘即可得到 1998 ~ 2002 年期间的 $(1 + \text{门店数量 4 年累积增长率})$ 的值。 $(1 + \text{门店数量 4 年累积增长率})$ 的值也可以由最终门店数量 4 036 除以最初门店数量 5 382 得到, 两种计算方法得到的结果是相同的, 即

$$\begin{aligned}\frac{4\,036}{5\,382} &= \left(\frac{5\,023}{5\,382}\right) \left(\frac{5\,129}{5\,023}\right) \left(\frac{4\,614}{5\,129}\right) \left(\frac{4\,036}{4\,614}\right) \\ &= (1 + g_1)(1 + g_2)(1 + g_3)(1 + g_4) \\ 0.749\,907 &= (0.933\,296)(1.021\,103)(0.899\,591)(0.874\,729)\end{aligned}$$

等式的右边是历年 $(1 + \text{门店数量年增长率})$ 的连乘积。回忆一下, 利用式 (1-14), 我们对 $\frac{4\,036}{5\,382} = 0.749\,907$ 求了 4 次方根。实际上, 我们所做的是在求解 g 的值, 该 g 值经过 4 次复利后 (即 $(1 + g)^4$) 等于历年 $(1 + \text{门店数量年增长率})$ 的连乘积。[⊖]

总之, 我们不必像表 1-4 那样通过计算出所有期间的增长率来求得复利增长率 g 。然而有时, 期间增长率更有趣或者能够提供给我们更多的信息。例如, 在 2000 年的 1 年中, 布兰兹有限公司增加了其门店数量。通过表 1-4 的计算, 我们还可以分析其增长率的变动情况。那么, 布兰兹有限公司是如何在此期间, 在不增加运营门店的情况下, 保持其销售收入基本不变的呢? 布兰兹有限公司在信息披露中提及在此期间内公司每平方英尺的销售额上升了。

复利增长率是一个极好的度量多期增长情况的汇总指标。在布兰兹有限公司的例子中, 复利增长率 -6.9% 是一个简单增长率。当它与 1 相加, 复利 4 年, 再乘以 1998 年运营门店的数量, 我们就可以得到 2002 年运营门店的数量。

1.7.2 求解期数

在这一节中, 我们将向你展示如何在给定现值、将来值和利率或增长率的情况下, 来对于期数进行求解。

例 1-19 一项投资达到某个特定值所需的复利期数

你想要确定一项 10 000 000 欧元的投资要多久才能价值翻倍。当前的利率为 7%, 按年度复利。那么要多久才能使得 10 000 000 欧元翻倍成 20 000 000 欧元?

解: 我们利用式 (1-2), $FV_N = PV(1 + r)^N$, 来求解期数 N 。具体过程如下:

$$\begin{aligned}(1 + r)^N &= FV_N / PV = 2 \\ N \ln(1 + r) &= \ln(2) \\ N &= \ln(2) / \ln(1 + r) \\ &= \ln(2) / \ln(1.07) = 10.24\end{aligned}$$

⊖ 我们在这里所计算的复利增长率实际上是计算几何平均值的一个例子, 具体来讲它是增长率的几何平均值。我们将在介绍统计概念的章节中给出几何平均的定义。

故在 7% 的利率水平下, 要经过大约 10 年才能使得最初的 10 000 000 欧元投资增长到 20 000 000 欧元。在求解 N 的表达式 $(1.07)^N = 2.0$ 时需要对等式两边取自然对数, 并利用对数性质 $\ln(x^N) = N\ln(x)$ 。一般地, 我们可以求得 $N = [\ln(FV/PV)]/\ln(1+r)$ 。这里, $N = [\ln(€20\,000\,000/€10\,000\,000)]/\ln(1.07) = \ln(2)/\ln(1.07) = 10.24$ 。[⊖]

1.7.3 求解年金支付额

在这一节中, 我们将讨论如何来求解年金支付额。抵押贷款、汽车消费贷款和退休储蓄计划都是年金公式应用的经典例子。

例 1-20 月度复利计息时达到将来值所需的年金支付额

你正计划购买一套价值 120 000 美元的住房。你会首付 20 000 美元, 其余部分将通过一份按月度复利计息的 30 年的固定利率抵押贷款来分期偿还。首笔支付将在 $t=1$ 时刻发生。当前抵押贷款的利率报价为 8%, 按月度复利。那么, 你每月应偿还的抵押贷款额为多少?

解: 银行将会确定抵押贷款的支付额, 以使得在期间报价利率的水平下, 所有支付额的现值等于贷款额 (在这个例子中为 100 000 美元)。在记住这一事实之后, 我们可以利用式

(1-11), $PV = A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right]$, 来求解年金支付额 A , 即通过将现值除以现值年金因子:

$$PV = \$100\,000$$

$$r_s = 8\% = 0.08$$

$$m = 12$$

$$r_s/m = 0.08/12 = 0.006\,667$$

$$N = 30$$

$$mN = 12 \times 30 = 360$$

$$\text{现值年金因子} = \frac{1 - \frac{1}{[1 + (r_s/m)]^{mN}}}{r_s/m} = \frac{1 - \frac{1}{(0.006\,667)^{360}}}{0.006\,667} = 136.283\,494$$

$$A = \frac{PV}{\text{现值年金因子}} = \frac{\$100\,000}{136.283\,494} = \$733.76$$

100 000 美元的贷款数额等价于一个 360 个月每月支付 733.76 美元并以 8% 报价年利率计息的年金。抵押贷款问题是求解年金支付额的一个相对直接的应用。

下面我们将转向退休金计划问题。这个问题描述了这样一个复杂的情况, 即一个人想要在退休时获得一笔特定的退休金收入。在整个生命周期中, 这个人可能只能在早期存入较小数额的存款, 但是在之后的人生中, 他可能会拥有一定的金融资源以存入更多的金额。储蓄计划常常涉及

⊖ 为了能快速估计期数, 实务工作者有时会使用一种特别的规则, 称为 72 规则: 即用 72 去除以报价利率来得到近似的使初始投资翻倍所需的年数。在这个例子中, 所估计的年数为 $72/7 = 10.3$ 年。72 规则是基于人们粗略的日常观察而得到的。人们发现在 6% 的利率水平下, 投资额翻倍一般需要 12 年, 于是就有了 $6 \times 12 = 72$ 。而在 3% 的利率水平下, 我们会猜测它将经过两倍的年数使投资翻倍, 即有 $3 \times 24 = 72$ 。

不等额的现金流，我们会在本章的最后部分讨论这个问题。当处理不等额现金流时，我们会最大限度地利用现金流可加性原理（cash flow additivity principle），即在相同时点上的美元金额是可加的。

例 1-21 未来年金流入所需筹集的投影年金金额

吉尔·格兰特（Jill Grant）今年 22 岁（在 $t=0$ 时刻），她现在正在为她 63 岁的退休进行打算（在 $t=41$ 时刻）。她计划在接下来的 15 年每年存入 2 000 美元（ $t=1$ 时刻至 $t=15$ 时刻）。她想要有 20 期的每年 100 000 美元的退休金收入，其首笔退休金支付始于 $t=41$ 时刻。那么，在 $t=16$ 时刻至 $t=40$ 时刻之间，每年格兰特要储蓄多少金额才能达到她所制定的退休金目标？假设她计划投资于平均年回报率 8% 的充分分散化的股票与债券型共同基金之中。

解：为了帮助解决该问题，我们将信息标注在一条时间轴上。如图 1-8 所示，格兰特将在第 1~15 年间每年存入 2 000 美元（现金流出）。从第 41 年开始，格兰特将开始在接下去的 20 年内每年提取退休金收入 100 000 美元。在时间轴上，年储蓄额用括号中的数值（2 美元）来表示这是一笔现金流出。我们的问题是求解从第 16~40 年的储蓄额，我们将其记为 X 。

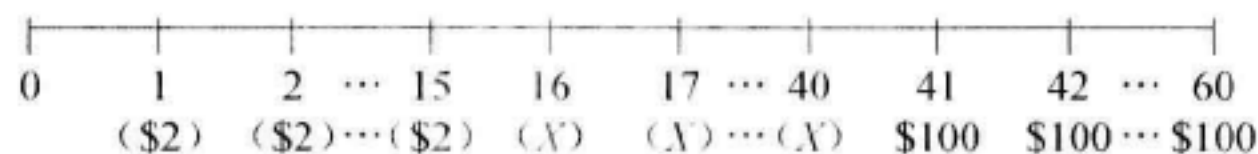


图 1-8 求解缺失的年金支付额（以千计）

该问题的求解需要利用如下的关系：储蓄的现值（流出）等于退休金收入的现值（流入）。我们可以将所有的美元数额换算到 $t=40$ 时刻或者 $t=15$ 时刻，并由此求解 X 。

我们不妨在 $t=15$ 时刻对所有的美元金额进行估值计算（我们鼓励读者在 $t=40$ 时刻估值的基础上重新将该问题再做一遍）。在 $t=15$ 时刻，首笔储蓄额 X 将在 1 期之后支付（在 $t=16$ 时刻）。因此，我们可以使用普通年金的现值公式对一个大小为 X 的现金流序列进行换算。

这个问题涉及三个等额现金流序列。我们基本的想法是退休金收入的现值必须与格兰特之前的储蓄额的现值相等。于是，我们的求解步骤如下：

1. 求解每年 2 000 美元储蓄在 $t=15$ 时刻的将来值。该值告诉我们格兰特在 $t=15$ 时刻将已经存储了多少金额。
2. 求解退休金收入在 $t=15$ 时刻的现值。该值告诉我们格兰特所要求的退休金目标金额是多少（同样是在 $t=15$ 时刻）。这里需要两个子步骤。首先，在 $t=40$ 时刻计算每年支付 100 000 美元年金的现值。这里我们利用了年金的现值公式。（注意到这里的现值是在 $t=40$ 时刻的，因为首期支付是在 $t=41$ 时刻发生的。）其次，将该现值折现到 $t=15$ 时刻（一共是 25 个期间）。
3. 现在我们可以计算格兰特的储蓄额（第一步所得的结果）和她所要求的退休金目标金额（第二步所得的结果）之差。她在 $t=16$ 时刻至 $t=40$ 时刻间的储蓄额的现值必须与其储蓄将来值与退休金收入的现值之差相等。

我们的目标是确定格兰特在 $t=16$ 时刻至 $t=40$ 时刻的 25 年间每年的储蓄额。我们从计算每年 2 000 美元的储蓄在 $t=15$ 时刻的将来值开始，具体计算如下：

$$A = \$2\,000$$

$$r = 8\% = 0.08$$

$$N = 15$$

$$\begin{aligned} FV &= A \left[\frac{(1+r)^N - 1}{r} \right] \\ &= \$2\,000 \left[\frac{(1.08)^{15} - 1}{0.08} \right] \\ &= \$2\,000(27.152\,114) \\ &= \$54\,304.23 \end{aligned}$$

在 $t=15$ 时刻，格兰特最初的储蓄将增长到 54 304.23 美元。

现在我们需要知道格兰特所需的退休金收入在 $t=15$ 时刻的价值。正如之前所述，计算退休金的现值需要两个子步骤。首先，利用式 (1-11) 求解在 $t=40$ 时刻的现值；其次，将该现值折现到 $t=15$ 时刻。现在我们求解 $t=40$ 时刻退休金收入的现值：

$$A = \$100\,000$$

$$r = 8\% = 0.08$$

$$N = 20$$

$$\begin{aligned} PV &= A \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \\ &= \$100\,000 \left[\frac{1 - \frac{1}{(1.08)^{20}}}{0.08} \right] \\ &= \$100\,000(9.818\,147) \\ &= \$981\,814.74 \end{aligned}$$

由于这是在 $t=40$ 时刻的现值，因此我们现在必须将其一次性折现到 $t=15$ 时刻：

$$FV_N = \$981\,814.74$$

$$N = 25$$

$$r = 8\% = 0.08$$

$$\begin{aligned} PV &= FV_N(1+r)^{-N} \\ &= \$981\,814.74(1.08)^{-25} \\ &= \$981\,814.74(0.146\,018) \\ &= \$143\,362.53 \end{aligned}$$

现在回忆一下，之前我们算得格兰特在 $t=15$ 时刻将已经储蓄 54 304.23 美元。因此，从现值的角度来看， $t=16$ 时刻至 $t=40$ 时刻的年金必须与已经储蓄的金额（54 304.23 美元）和所需要的退休金数量（143 362.53 美元）之间的差额相等。该值等于 $\$143\,362.53 - \$54\,304.23 = \$89\,058.30$ 。因此，我们现在必须求解从 $t=16$ 时刻至 $t=40$ 时刻的现值为 89 058.30 美元的年金的每期支付额 A 。我们计算该年金支付额的过程如下：

$$\begin{aligned}PV &= \$89\,058.30 \\r &= 8\% = 0.08 \\N &= 25\end{aligned}$$
$$\begin{aligned}\text{现值年金因子} &= \left[\frac{1 - \frac{1}{(1+r)^N}}{r} \right] \\&= \left[\frac{1 - \frac{1}{(1.08)^{25}}}{0.08} \right] \\&= 10.674\,776 \\A &= PV / \text{现值年金因子} \\&= \$89\,058.30 / 10.674\,776 \\&= \$8\,342.87\end{aligned}$$

格兰特需要在 $t=16$ 时刻至 $t=40$ 时刻间将每年的储蓄额增加到 8 342.87 美元，以达到她退休的目标，即在 $t=40$ 时刻最后一笔储蓄支付之后她将获得一笔价值 981 814.74 美元的资金。

1.7.4 现值和将来值换算关系的回顾

正如我们已经讨论的，求解现值和将来值涉及将货币金额在时间轴上不同时间点间的移动。该操作之所以是可能的，是因为现值和将来值是在时间上分离的等价度量指标。表 1-5 举例说明了这一等价性；它列示了 5 笔现金流发生的时间，它们的现值是在 $t=0$ 时刻计算的，而它们的将来值是在 $t=5$ 时刻计算的。

表 1-5 现值与未来值的等价性 (美元)			
时点	现金流	在 $t=0$ 时刻的现值	在 $t=5$ 时刻的将来值
1	1 000	$1\,000 (1.05)^{-1} = 952.38$	$1\,000 (1.05)^4 = 1\,215.51$
2	1 000	$1\,000 (1.05)^{-2} = 907.03$	$1\,000 (1.05)^3 = 1\,157.63$
3	1 000	$1\,000 (1.05)^{-3} = 863.84$	$1\,000 (1.05)^2 = 1\,102.50$
4	1 000	$1\,000 (1.05)^{-4} = 822.70$	$1\,000 (1.05)^1 = 1\,050.00$
5	1 000	$1\,000 (1.05)^{-5} = 783.53$	$1\,000 (1.05)^0 = 1\,000.00$
		合计: 4 329.48	合计: 5 525.64

要解释表 1-5，我们首先从第三列开始，该列的数字代表现值。注意到每笔 1 000 美元的现金支付被以适当的期数折现为 $t=0$ 时刻的现值。现值 4 329.48 美元正好等价于这一现金流序列。这一信息揭示了一个重要的观点：一次性支付的现金流可以产生一个年金。如果我们将这个一次性支付放在一个账户之中，并让其在整个期间中以报价利率计息，那么我们将得到一个与之相等价的年金。分期付款如抵押贷款和汽车消费贷款，都是应用这一原理的例子。

为了了解一笔一次性支付现金流如何产生一个年金，假设在今天我们以 5% 利率在银行中存入 4 329.48 美元。我们可以通过式 (1-11) 计算出年金支付额的大小。通过求解 A ，我们得到：

$$A = \frac{PV}{\frac{1 - [1/(1+r)^N]}{r}}$$

$$\begin{aligned} &= \frac{\$4\,329.48}{\frac{1 - [1/(1.05)^5]}{0.05}} \\ &= \$1\,000 \end{aligned}$$

表 1-6 向我们显示最初的 4 329.48 美元投资是如何在接下来的 5 年内实际产生 5 个 1 000 美元的提款额的。

表 1-6 最初现值如何产生一个年金 (美元)				
期间	期初所拥有的金额	期末提款前所拥有的金额	提款额	提款后所拥有的金额
1	4 329.48	4 329.48 (1.05) = 4 545.95	1 000	3 545.95
2	3 545.95	3 545.95 (1.05) = 3 723.25	1 000	2 723.25
3	2 723.25	2 723.25 (1.05) = 2 859.41	1 000	1 859.41
4	1 859.41	1 859.41 (1.05) = 1 952.38	1 000	952.38
5	952.38	952.38 (1.05) = 1 000	1 000	0

要解释表 1-6，我们先从 $t=0$ 时刻最初的现值 4 329.48 美元开始。从 $t=0$ 时刻到 $t=1$ 时刻，初始投资获得了 5% 的利息，产生将来值 $\$4\,329.48 (1.05) = \$4\,545.95$ 。然后我们从账户中取出 1 000 美元，剩下 $\$4\,545.95 - \$1\,000 = \$3\,545.95$ （该数字在表中第一个期间的最后一列中给出）。在下一个期间，我们同样获得了 1 年的利息并从中提款 1 000 美元。在第 4 次提款之后，账户余额为 952.38 美元，它将继续获得 5% 的利息。于是，这一金额将在这一年中增长到 1 000 美元，而这正好满足我们最后一次的提款额度。因此最初的现值，在以 5% 的利率投资 5 年后，将产生每年 1 000 美元的 5 年期普通年金。该初始投资的现值正好等价于这个年金。

现在我们可以来看一下，将来值是如何与年金相关的。在表 1-5 中，我们所报告的年金将来值为 5 525.64 美元。我们通过将最初的 1 000 美元支付向前复利 4 期，将第二期的 1 000 美元支付向前复利 3 期，以此类推。我们可以将所得到的 5 个在 $t=5$ 时刻的将来值加总。于是，年金与 $t=5$ 时刻的 5 525.64 美元和 $t=0$ 时刻的 4 329.48 美元等价。因此，这两个美元数值是等价的。我们可以通过求解将来值 5 525.64 美元的现值（ $\$5\,525.64 (1.05^{-5}) = \$4\,329.48$ ），来证明这两者间的等价性。在我们揭示一次性支付可以产生一个年金的同时，我们发现了如上的结果。

总结一下至今我们所学到的内容：一次性支付额可以视为等价于一个年金，而年金可以认为等价于其将来值。因此，只要标记在同一时点上，现值、将来值和现金流序列都可以被认为是相互等价的。

1.7.5 现金流可加性原理

现金流可加性原理（指标记于相同时点的货币数量可以相加的原理）是有关货币时间价值的数学计算中最为重要的概念之一。我们在之前的讨论中已经提及并使用了该原理，本节将提供一个关于该原理的例子作为参考。

考虑两个现金流序列，它们分别标示在图 1-9 中的时间轴上。这两个现金流序列分别记为 A 和 B。如果我们假设年利率是 2%，我们可以求得每个现金流序列的将来值。现金流序列 A 的将来值为 $\$100 (1.02) + \$100 = \$202$ ，现金流序列 B 的将来值为 $\$200 (1.02) + \$200 = \$404$ 。通过使用我们到目前为止所学到的方法，现金流序列 (A + B) 的将来值为 $\$202 + \$404 = \$606$ 。另外一种计算该将来值的方法是将这两个现金流序列 A 和 B 中的每个时点的现金流分别加总（称为

A + B)，然后去求解这样一个合并现金流序列的将来值，如图 1-9 所示。

图 1-9 中的第 3 条时间轴反映的就是合并的现金流序列。序列 A 在 $t=1$ 时刻有一笔 100 美元的现金流，而序列 B 在 $t=1$ 时刻有一笔 200 美元的现金流。因此，合并现金流序列在 $t=1$ 就有一笔 300 美元的现金流。我们同样可以计算出在 $t=2$ 时刻合并现金流序列的现金流。于是，合并现金流序列（A + B）的将来值为 $\$300(1.02) + \$300 = \$606$ 。求得的这一结果与我们之前将两个现金流序列将来值加总得到的结果是相同的。

可加性和等价性原理同样也出现在其他常见的情况之中。假如一个现金流序列在第 1 年年末支付 4 美元，在第 2 年年末支付 24 美元（实际上是分别支付 4 美元和 20 美元）。如果我们不想通过分别求解第 1 年支付 4 美元和第 2 年支付 20 美元的现值来获得该现金流序列的现值，那么我们可以将该现金流序列视为一个每期支付 4 美元的普通年金加上一个第 2 年年末一次性支付 20 美元的现金流。如果贴现率为 6%，那么每期支付 4 美元的普通年金的现值为 7.33 美元，而一次性支付 20 美元的现值为 17.80 美元。它们的总和为 25.13 美元，此即为该现金流序列的现值。

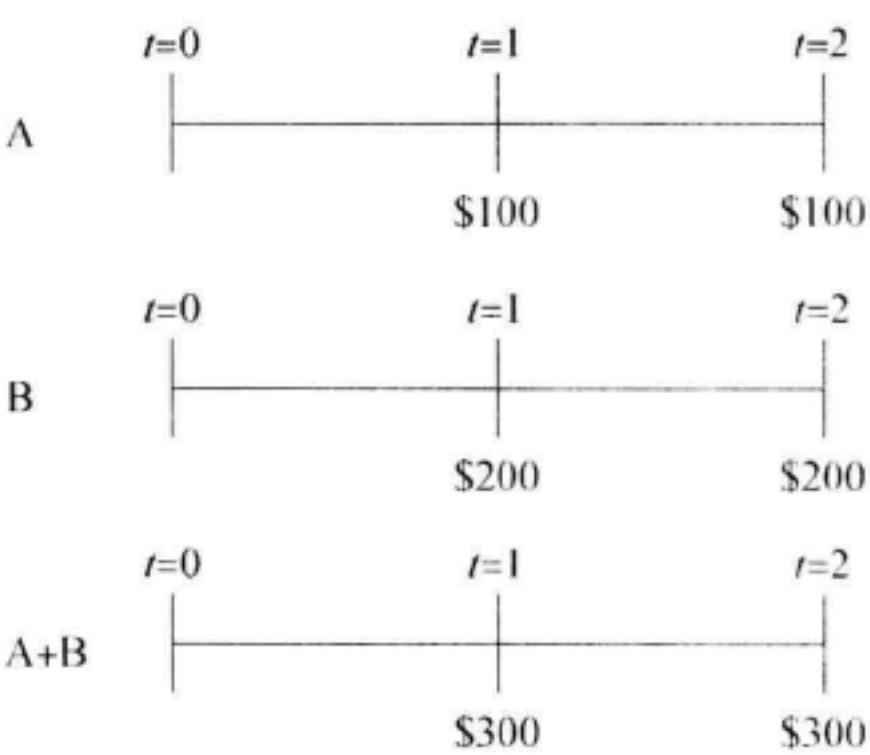


图 1-9 两个现金流序列的可加性



贴现现金流的应用

2.1 引言

作为投资分析师，我们的许多工作都涉及对具有当前或者未来现金流的交易进行评估。在货币的时间价值（the time value of money, TVM）一章中，我们介绍了解决此类问题所需的数学工具，并且举例说明了解决其中若干主要类型的问题所需要的一些计算技巧。在本章中我们将把目光转向具体的应用问题。当投资分析师对权益类证券、固定收益证券和衍生品中的某个问题进行研究时，必须掌握 TVM 或贴现现金流分析在权益类证券、固定收益证券和衍生品的分析中的许多应用技巧。在本章中，我们将选择一些重要的有关货币时间价值的应用问题加以介绍：利用净现值和内部收益率等工具来评估现金流序列的价值、度量投资组合收益率以及计算货币市场收益率。这些应用问题本身是十分重要的，并且在解决这些应用问题的过程中所引入的一些概念还将在许多其他投资分析问题上反复出现。

本章的内容安排如下：第 2 节介绍两个重要的有关货币时间价值的概念——净现值和内部收益率。在这两个概念的基础上，第 3 节将讨论投资管理中的一个重要问题——投资组合收益率的度量。投资经理经常要面对将资金进行短期投资的任务，为了了解各种投资机会，他们必须熟悉货币市场收益率的计算方式，因此，本章将在最后的第 4 节中讨论这一问题。

2.2 净现值和内部收益率

在任何金融领域应用贴现现金流分析的过程中，我们都会不断地遇到两个概念：净现值和内部收益率。在下面的几节中，我们将对这两个基本概念进行介绍。

因为净现值和内部收益率的应用范围覆盖整个金融领域，所以我们可以许多情形下对这两个概念进行讨论。而资本预算可以作为我们讨论的一个有代表性的起点。因为股票和固定收益证券的分析师都必须能够对公司管理者运用公司资产进行投资的业绩好坏进行评估，所以资本预算不仅在公司财务中十分重要，在证券分析中也发挥着关键性的作用。在大部分的商业活动中，主要有三个涉及金融决策的领域。资本预算（capital budgeting）是指将资金配置到一些相对长期的项目或投资

机会之中。从资本预算的角度来看,公司就是一个由项目和投资构成的资产组合。资本结构 (capital structure) 是指公司对于将要进行的投资项目所选择的长期融资方式。营运资本管理 (working capital management) 是公司对于短期资产 (如存货) 和短期负债 (如欠供应商的款项) 的管理。

2.2.1 净现值和净现值准则

净现值给出了一种评估某项投资的价值的方式,而净现值准则是用来在不同投资项目中进行选择的一种方法。一项投资的净现值 (net present value) 等于其所有现金流入的现值减去所有现金流出的现值。净现值中的“净”字指的是将投资中的现金流出 (成本) 的现值从其现金流入 (收益) 的现值中减去,从而得到净的收益。

计算净现值和应用净现值准则的步骤如下:

1. 识别出与投资项目相关的所有现金流——所有流入和流出的现金流。[⊖]
2. 为投资项目确定适当的贴现率或机会成本 r 。[⊖]
3. 使用第二步中所确定的贴现率,计算出每笔现金流的现值。(现金流入的符号为正,它会增加净现值;而现金流出的符号为负,它会减少净现值。)
4. 将所有求得的现值加总。所有现金流 (包括流入和流出) 现值的总和即为投资项目的净现值。
5. 应用净现值准则 (NPV rule): 若投资项目的净现值为正,则投资者应该进行投资;若投资项目的净现值为负,则投资者就不应该进行该项投资。如果投资者要在两个候选投资项目中选择其中一个进行投资 (即两个互斥的投资项目),那么投资者应该选择那个具有较高的正净现值的候选项目进行投资。

净现值准则的含义是什么?在计算一个投资项目的净现值时,我们用资本的机会成本的估计值作为贴现率。资本的机会成本是指投资者为进行该项投资而放弃的其他可供选择的收益。当净现值为正时,该项投资就会带来价值,因为该项投资的收益已经超出了进行该项投资所需承担的资本的机会成本。因此,公司接受具有正净现值的投资项目会增加其股东的财富。而对于个人投资者来说,正净现值的投资项目同样也会增加个人的财富;反之,负净现值的投资项目会减少个人的财富。

当我们使用净现值准则来解决问题时,下面的公式会对我们有所帮助:

$$NPV = \sum_{t=0}^N \frac{CF_t}{(1+r)^t} \tag{2-1}$$

式中, CF_t 表示在 t 时刻的期望净现金流;

N 表示投资项目的期限;

r 表示贴现率或资本的机会成本。

通常,我们所表述的各个输入变量应具有一致的时间基准:如果现金流是按年算的,那么 N 就应该是该项目的投资年数,而 r 就应该为年利率。例如,假设你现在要考察一个初始投资 200

⊖ 在进行现金流估计时,我们遵照两条规则。首先,这里的现金流只包括由投资项目所产生的增量现金流 (incremental cash flows);我们不将沉没成本 (在项目投资之前已经发生的成本) 包括在内。其次,我们使用税后现金流来对税收效应进行调整。对于资本预算的这些内容以及其他问题的详细讨论,读者可以参见 Brealey 和 Myer (2003)。

⊖ 加权平均资本成本 (Weighted-Average Cost of Capital, WACC) 经常被用来对现金流进行贴现。该值是公司普通股、优先股和长期债务税后要求收益率的加权平均值,其权重为公司目标资本结构中各个融资来源所占的比重。关于资本成本相关内容的详细讨论,读者可以参见 Brealey 和 Myer(2003)。

万美元 ($CF_0 = -\$2\,000\,000$) 的投资方案。你预期该投资项目将在未来三年产生一系列正的净现金流, 其值分别为第1年年末 $CF_1 = \$500\,000$ 、第2年年末 $CF_2 = \$750\,000$ 、第3年年末 $CF_3 = \$1\,350\,000$ 。假设我们使用10%的贴现率, 那么该项目的净现值就可以计算如下:

$$\begin{aligned} NPV &= -\$2 + \frac{\$0.50}{1.10} + \frac{\$0.75}{1.10^2} + \frac{\$1.35}{1.10^3} \\ &= -\$2 + \$0.454\,545 + \$0.619\,835 + \$1.014\,275 \\ &= \$0.088\,655 (\text{百万}) \end{aligned}$$

由于净现值88 655美元为正, 因此, 根据净现值准则, 你应该接受该投资方案。

下面考虑一个利用净现值准则来评估研发项目的例子。

例 2-1 利用净现值准则来评估研发项目

作为 RAD 公司的一位投资分析师, 你正在评估当年的一项研发项目 (R&D)。管理层已经宣布他们打算在该研发项目中投入100万美元。由该项目所带来的增量净现金流预计为每年150 000美元的永续年金。RAD公司的资本机会成本为10%。

1. 根据净现值准则来判断 RAD 公司的研发项目是否会给股东带来收益。
2. 如果 RAD 公司的资本机会成本是15%而非10%, 那么你在上一题中的结论是否会发生改变?

解:

1. 150 000美元的持续净现金流序列形成了一个永续年金, 我们可以将其记为 $\overline{CF}$ 。永续年金的现值为 $\overline{CF}/r$, 所以我们可以计算出该项目的净现值为

$$NPV = CF_0 + \frac{\overline{CF}}{r} = -\$1\,000\,000 + \frac{\$150\,000}{0.10} = \$500\,000$$

在机会成本为10%的情况下, 该项目现金流入的现值为150万美元。该项目的成本为即刻流出的100万美元; 因此, 它的净现值为500 000美元。由于净现值为正, 从而得出结论: RAD公司的研发项目将会给股东带来收益。

2. 在机会成本为15%的情况下, 可以按上述相同的方法计算净现值, 但这次要使用的贴现率为15%:

$$NPV = -\$1\,000\,000 + \frac{\$150\,000}{0.15} = \$0$$

在资本的机会成本较高时, 流入现金流的现值变小了, 从而该项目的净现值也变小了: 在15%的贴现率下, 该项目的净现值正好等于0。在 $NPV=0$ 的情形下, 该项目所产生的流入现金流只能刚好弥补股东进行该项投资的机会成本。因此, 当公司投资于零净现值的项目时, 虽然公司的规模变大了, 但是股东的财富并没有增加。

2.2.2 内部收益率和内部收益率准则

财务经理经常想用一个单一的数值指标来表示由投资项目所产生的收益率。在投资应用中 (包括资本预算中), 一个最常使用的收益率的计算量为内部收益率。内部收益率准则是用来对不同投资方案进行选择的第二种方法。内部收益率 (internal rate of return, IRR) 是使得项目净现值为零的贴现率。该贴现率使得投资成本 (现金流出) 的现值与投资收益 (现金流入) 的现值相

等。这一收益率之所以被称为是“内部”的，是因为它的计算只依赖于投资项目的现金流，并不需要其他外部数据。因此，我们可以将 IRR 的概念应用到任何可以由一系列现金流表示的投资项目中去。在对于债券的研究中，我们会遇到 IRR 的概念，在那里它被称为到期收益率。在本章的后面部分，我们将把 IRR 作为投资组合的货币加权收益率来进行探讨。

然而，在继续讨论之前，我们必须补充一个关于 IRR 含义的重要提示：即使我们的现金流贴现计算都是正确的，我们也只有在可以对所有期间现金流以内部收益率进行再投资的情况下，才能在整个投资期限上实现等于内部收益率的复利收益率。假设一个项目的 IRR 为 15%，但是我们持续地将由这个项目所产生的现金流以低于 15% 的利率进行再投资，那么在这种情况下，我们将只能获得低于 15% 的投资收益。（如果我们以高于 15% 的水平进行再投资，那么这一原则将会对我们有利。）

让我们回到内部收益率的定义，我们可以将其以数学式表述为：

$$NPV = CF_0 + \frac{CF_1}{(1 + IRR)^1} + \frac{CF_2}{(1 + IRR)^2} + \cdots + \frac{CF_N}{(1 + IRR)^N} = 0 \quad (2-2)$$

同样地，式（2-2）中的 IRR 必须与现金流的时间期限相匹配。如果现金流是季度的，我们就可以由式（2-2）获得季度的内部收益率。我们可以接着将其化为年度的内部收益率。对于一些简单的项目，在 0 时点的现金流 CF_0 为单笔资本支出或初始投资；而在 0 时点之后的现金流为投资项目的正收益。在这种情况下，我们就认为 $CF_0 = -\text{投资额}$ （负号代表现金流出）。于是，我们可以重新改写式（2-2）为如下的形式（这一形式在上述情况下是很有用的）：

$$\text{投资额} = \frac{CF_1}{(1 + IRR)^1} + \frac{CF_2}{(1 + IRR)^2} + \cdots + \frac{CF_N}{(1 + IRR)^N}$$

对于大部分现实问题来说，金融分析师会使用软件、电子表格或者金融计算器去求解内部收益率的方程，因此你应该熟悉这些计算工具。^②

使用 IRR 的投资决策准则，即内部收益率准则（IRR rule），可以表述为：“接受那些 IRR 高于资本机会成本的项目或者投资。” IRR 准则使用资本的机会成本作为一个门槛比率（hurdle rate），或者是项目被接受时其内部收益率所必须超过的水平。注意到如果资本的机会成本等于内部收益率，那么该项目的净现值将等于 0。如果资本的机会成本小于内部收益率，那么其净现值将大于 0（使用低于内部收益率的贴现率将使得净现值为正）。记住这些要点，我们就可以来解决下面例题中两个与内部收益率有关的问题。

例 2-2 利用内部收益率准则来评估研发项目

在之前的 RAD 公司的例子中，最初投资为 100 万美元，而项目的现金流为 \$150 000 的永续年金。现在你可能会关心该项目的内部收益率。现在我们要解决下列问题：

- （1）写出决定该研发项目的内部收益率的方程式。
- （2）计算 IRR。

② 在一些资本预算的现实问题中，在初始投资（带有负号）之后可能会紧跟着一系列的现金流入（带有正号）和现金流出（带有负号）。在这些情况下，投资项目可能会有不止一个内部收益率。可能出现多个解是内部收益率理论的一个局限。

解:

(1) 求解内部收益率等价于寻找使得净现值等于0时的贴现率。因为项目的现金流是永续年金,所以可以将净现值的等式写为:

$$NPV = -\text{投资} + \frac{\overline{CF}}{IRR} = 0$$

$$NPV = -\$1\,000\,000 + \frac{\$150\,000}{IRR} = 0$$

或者写成

$$\text{投资} = \frac{\overline{CF}}{IRR}$$

$$\$1\,000\,000 = \frac{\$150\,000}{IRR}$$

(2) 我们可以求解内部收益率为 $IRR = \frac{\$150\,000}{\$1\,000\,000} = 0.15$ 或者 15%。15%的这个解与

内部收益率的定义是一致的。在例2-1中,你曾经发现15%的贴现率会使得项目的净现值等于0。因此,根据定义,该项目的内部收益率一定是15%。如果资本的机会成本同样是15%,那么该研发项目的收益正好弥补了其机会成本,并且没有增加也没有减少股东财富。如果资本的机会成本低于15%,那么内部收益率准则表明管理层应该投资于该项目,因为它的收益超过了机会成本。如果机会成本高于15%,那么内部收益率准则告诉管理层应该拒绝该项目。因此,在给定机会成本的条件下,内部收益率准则和净现值准则在这个例子中将得到相同的结论。

例2-3 同时考虑内部收益率和净现值准则

日本的影山(Kageyama)有限公司正考虑是否要开设一家生产用于手机中的电容器的新工厂。该工厂将需要一笔¥1 000 000 000的投资。该工厂预计将在接下去的5年每年产生¥294 800 000的等额现金流。根据其财务报告中的信息,影山公司对于该类项目的资本机会成本设定为11%。

(1) 根据净现值准则确定该项目是否对影山公司的股东有利。

(2) 根据内部收益率准则确定该项目是否对影山公司的股东有利。

解:

(1) 现金流可以被分为初始的1 000 000 000日元现金流出和五年期每年294 800 000日元现金流入的普通年金。年金的现值表示为 $A[1 - (1 + r)^{-N}]/r$, 其中A是等额的年金支付额。因此,用百万日元作为单位,其数值为

$$\begin{aligned} NPV &= -1\,000 + 294.8[1 - (1.11)^{-5}]/0.11 \\ &= -1\,000 + 1\,089.55 = 89.55。 \end{aligned}$$

因为该项目的净现值为正的89.55万日元,所以它将使得影山公司的股东获益。

(2) 该项目的内部收益率是如下等式的解:

$$NPV = -1\,000 + 294.8[1 - (1 + IRR)^{-5}]/IRR = 0$$

该项目具有正的净现值告诉我们内部收益率一定大于11%。使用计算器，我们可以求得IRR为0.145 012 或者是14.50%。表2-1给出了大部分金融计算器的按键。

因为内部收益率14.50%比该项目的机会成本要大，所以该项目将会给影山公司的股东带来收益。无论使用内部收益率准则还是使用净现值准则，影山公司都会做出相同的决策，即建设该工厂。

表 2-1 IRR 的计算

大部分计算器所用的符号	本例中问题对应的数值	大部分计算器所用的符号	本例中问题对应的数值
<i>N</i>	5	PMT	294.8
% <i>i</i> compute	<i>X</i>	FV	<i>n/a</i> (=0)
PV	-1 000		

在上面的例子中，价值的创造是明显的：对于一笔1 000 000 000 日元的投资，影山公司创造了一个价值为1 089 550 000 日元的项目，价值增长了89 550 000 日元。另外一个判断价值创造的方法是，将初始投资转化成一項对于项目产生的年经营现金流的资本要求。回忆一下，该项目产生的年经营现金流为294 800 000 日元。如果我们减去270 570 310 日元的资本要求（这是以11%利率折现，现值为1 000 000 000 日元的5年期年金的年支付额），我们会得到¥294 800 000 - ¥270 570 310 = ¥24 229 690。24 229 690 日元这一数值代表在考虑机会成本之后，未来5年每年的盈利。每年24 229 690 日元的5年期年金以11%的资本成本折现所得的现值正好是我们之前计算得到的项目的净现值：89 550 000 日元。因此，我们也可以通过将初始投资转化为年资本要求的现金流来计算净现值。

2.2.3 与内部收益率准则相关的问题

当项目间相互独立时（即投资于某一个项目的决策不影响是否投资于另一个项目的决策时），内部收益率准则和净现值准则给出了相同的（接受或拒绝的）决策。而当一家公司无法为其所有想要投资的项目融资时，即当项目间是互斥时，它就必须对所有的项目按照盈利能力由大到小进行排序。然而，根据内部收益率准则和净现值准则进行排序的结果可能并不一定相同。用内部收益率准则和净现值准则对项目进行排序，得到不同结果的情况有：

- 当项目间的投资规模不同时（投资规模是由所接受项目的所需投资额来度量的）。
- 当项目的现金流序列发生的时点不同时。

当内部收益率准则和净现值准则在对项目的排序产生冲突时，我们应该采纳净现值准则的结果。为什么优先参考净现值准则的结果？因为投资项目的净现值代表由投资项目给股东财富所带来的预期增值，而我们将股东财富最大化作为一个最基本的公司财务目标。为了举例说明选择净现值准则的原因，我们可以首先考虑不同投资规模项目的情况。假设一个公司只能用€30 000 来投资。[⊖]现在，该公司有两个一期的、如表2-2中A和B所示的投资项目可以选择。

表 2-2 不同规模的互斥投资项目的内部收益率和净现值

项目	在 <i>t</i> = 0 处的投资 (€)	在 <i>t</i> = 1 处的现金流 (€)	内部收益率 (%)	8%水平下的净现值 (€)
A	-10 000	15 000	50	3 889.89
B	-30 000	42 000	40	8 889.89

⊖ 或者假设两个项目需要相同的物理的或其他资源，因此只有一个项目可以被选择。

项目 A 要求即刻投资€10 000。该项目将在 $t = 1$ 处获得一笔€15 000 的现金支付。因为内部收益率是使得将来现金流的现值与投资成本相等的折现率，故该内部收益率等于 50%。如果我们假设资本的机会成本为 8%，那么项目 A 的净现值为€3 889.89。我们可以同样地计算出项目 B 的内部收益率和净现值分别为 40% 和€8 889.89。根据内部收益率准则和净现值准则，我们应该同时接受这两个项目，但是如果这么做的话，我们就需要€40 000（初始投资），这超出了我们所能获得的投资额，所以我们需要对这两个项目进行排序。那么根据内部收益率准则和净现值准则所获得的项目排序分别是怎样的呢？

内部收益率准则将项目 A 排在首位，因为它拥有较高的内部收益率。而净现值准则则将项目 B 排在首位，因为它具有较高的净现值。从而，净现值准则的结果与内部收益率准则的结果产生了矛盾（冲突），因为项目 A 具有较高的内部收益率而选择它，并不会最大限度地增加股东的财富。完全投资于项目 A 将会剩下€20 000 没有作任何投资。项目 A 将增加财富近€4 000，但是项目 B 将增加财富近€9 000。这两个项目的规模不同导致了两个准则进行排序时产生的结果不一致。

内部收益率和净现值也可能在现金流发生在不同时点的情况下，对于同样规模的项目排序不同。我们可以用表 2-3 中的项目 A 和 D 的例子说明这一情况。

表 2-3 对于现金流发生在不同时点的互斥项目的内部收益率和净现值

(金额单位：€)						
项目	CF ₀	CF ₁	CF ₂	CF ₃	内部收益率 (%)	8%水平下的净现值
A	-10 000	15 000	0	0	50.0	3 889.89
D	-10 000	0	0	21 220	28.5	6 845.12

CF₀、CF₁、CF₂、CF₃ 这四项分别表示在时间段 0、1、2、3 中所发生的现金流。项目 A 的内部收益率和上例中的结果一样。项目 D 的内部收益率按如下方式求解：

$$-10\,000 + 21\,220 / (1 + \text{IRR})^3 = 0$$

相比于项目 A 的 50%，项目 D 的内部收益率（IRR）为 28.5%。因为项目的现金流序列完全确定了内部收益率，因此内部收益率和按内部收益率准则的排序不受到外部利率或者贴现率的影响。内部收益率的计算进一步假设了我们能够以内部收益率进行再投资，所以我们一般不能将其解释为可获得的实际回报率。例如，对于项目 D，为了实现 28.5% 的收益，我们需要第一年在投资额€10 000 的基础上获取 28.5% 的收益，在第 2 年€10 000(1.285) = €12 850 的基础上也获得 28.5% 的收益，而在第三年€10 000(1.285)² = €16 512.25 的基础上同样也获得 28.5% 的收益。^①以 50% 或者 28.5% 的利率进行再投资可能不太现实。相反，计算净现值需要外部市场决定的贴现率，再投资将被认为是以这一贴现率进行的。根据净现值准则进行的排序依赖于外部贴现率的选择。这里，项目 D 具有更大但更远的现金流入（€21 220 对比€15 000）。因此，项目 D 在较低的贴现率^②下，具有比项目 A 更高的净现值。净现值准则对于再投资利率的假设更加现实并与经济更密切相关，因为它将市场确定的资本机会成本作为了贴现率。因此，净现值是从投资中获得的对股东财富的期望增加值。

总之，在处理互斥项目时，若内部收益率准则和净现值准则相互矛盾，我们会在其中选择净现值准则的结果。^③

① 期末的数值€10 000 (1.285)³ = €21 218，不同于表 2-3 中所列示€21 220，这是因为我们在计算内部收益率时进行了四舍五入。
② 存在一个交叉的贴现率，在该数值之上项目 A 的净现值高于项目 D。这一交叉贴现值为 18.94%。
③ 从技术上来说，对于再投资率的不同假设解释了内部收益率准则和净现值准则结果间的不同。内部收益率准则假设了公司所有投资现金流都获得内部收益率的回报，而净现值准则假设了现金流都以公司的资本机会成本进行再投资。净现值准则更加地现实。想要了解更多关于该内容和资本预算的相关主题，可以参见 Brealey 和 Myers (2003)

2.3 投资组合收益的度量

假设你是一个投资者，并且你想要评估你投资的成果。那么，你将面临两个相关但又不同的工作。第一个是业绩度量（performance measurement），即关于如何以逻辑和一致的方法来评估收益率。准确的业绩度量为你的第二项工作即业绩评估（performance appraisal）提供了基础。[Ⓐ]因此，业绩度量对于所有投资者和投资经理来说至关重要，因为它是所有进一步分析的基础。

在我们对于投资组合收益率度量的讨论中，我们将使用持有期回报率（holding period return, HPR）这个基本概念。持有期回报率指的是一个投资者在一个特定的持有期间中所获得的收益。对于一项在持有期期末产生一笔现金支付的投资来说，其持有期回报率为： $HPR = (P_1 - P_0 + D_1)/P_0$ ，其中 P_0 为初始投资， P_1 为在持有期期末所获得的价格，而 D_1 为在持有期期末由该投资所支付的现金。

特别地，当我们度量多期投资的业绩表现时，或者是当投资组合在持有期中既有资金投入又有资金取出时，投资组合业绩度量将成为一项具有挑战性的工作。两个可以用来度量的工具分别是货币加权收益率度量指标和时间加权收益率度量指标。我们讨论的第一个指标货币加权收益率，使用了我们在资本预算内容中所讨论过的一个概念，即内部收益率。

2.3.1 货币加权收益率

第一个我们要讨论的业绩度量的概念是内部收益率的计算。在投资管理的应用中，内部收益率常被称为货币加权收益率，因为它解释了投资组合中所有美元现金流流入和流出的时间和数量。[Ⓑ]

为了说明货币加权收益率，我们考虑一个两年期的投资。在时点 $t = 0$ ，一个投资者以 \$200 购买一股股票。在时点 $t = 1$ ，他再以 \$225 购买一股。在第 2 年年末， $t = 2$ ，他以 \$235 每股的价格将其两股卖出。在这两年中，该股票每年支付每股 \$5 的红利。 $t = 1$ 时的红利不做再投资。表 2-4 列出了所有现金流入与流出的情况。

表 2-4 现金流序列	
时间	支出
0	\$200 购买第一股
1	\$225 购买第二股
收入	
1	由第一股投资所获得 \$5 红利（不做再投资）
2	收到的 \$10 红利（\$5 每股，2 股）
2	在 \$235 卖出两股所获得的 \$470

这个投资组合的货币加权收益率是该投资组合两年的内部收益率。投资组合的内部收益率是使得现金流流出的现值减去现金流流入的现值等于 0 的利率 r 。或者我们有：

$$PV(\text{现金流出}) = PV(\text{现金流入})$$
$$200 + 225/(1 + r) = 5/(1 + r) + 480/(1 + r)^2$$

等式的左侧详细描述了现金流出的情况：在 $t = 0$ 时刻的 \$200 和在 $t = 1$ 时刻的 \$225。现金流出的 \$225 被折现到了一期之前，这是因为它发生在 $t = 1$ 时刻。等式的右侧详细描述了现金流流入的情况：在 $t = 1$ 时刻的 \$5（折现到一期以前）和在 $t = 2$ 时刻的 \$480（红利 \$10 加上 \$470 的

Ⓐ “业绩评估”这个术语是用来作为业绩评价的同义词。在之后的章节中我们将讨论一个业绩评价指标——夏普比率。
Ⓑ 在美国，货币加权收益率经常被称为美元加权收益率。我们沿袭了传统的货币加权收益率的标准表述来讨论内部收益率的概念。

卖出收益)。

为了求解货币加权收益率,我们或者使用一个能够使我们输入现金流序列的金融计算器,或者利用一个含有内部收益率函数的电子表格。^①第一步是按时间将净现金流序列归类。在这个例子中,我们在 $t=0$ 时刻有 $-\$200$ 的净现金流,在 $t=1$ 时刻有 $-\$200 = -\$225 + \$5$ 的净现金流,而在 $t=2$ 时刻有 $\$480$ 的净现金流。在将这些现金流序列输入之后,我们将利用电子表格或该计算器的内部收益率函数来求解货币加权收益率,其解为 9.39% 。^②

现在我们可以进一步仔细了解一下,该资产组合在这两年的每一年中都发生了些什么变化。在第一年,投资组合产生了一个一期的持有期回报率 $(\$5 + \$225 - \$200)/\$200 = 15\%$ 。在第2年年初,投资额为 $\$450$,由 $\$225$ (每股价格) $\times 2$ 股计算得到,因为 $\$5$ 红利支付后没有进行再投资。在第2年年末该投资组合清算后的收益为 $\$470$ (如表2-4中所详细描述)加上 $\$10$ 红利(也如表2-4中所显示的)。所以在第2年,投资组合所产生的持有期回报率为 $(\$10 + \$470 - \$450)/\$200 = 6.67\%$ 。平均的持有期收益为 $(15\% + 6.67\%)/2 = 10.84\%$ 。我们计算所得到的货币加权收益率 9.39% ,将更大的权重放在了第2年相对较差的表现上(6.67%),而不是第一年相对较好的表现上(15%)。这是因为第2年比第1年投资了更多的资金。因此,在这个意义上,以这种计算业绩表现回报率的方法被认为是“按货币加权的”。

作为一个评价投资经理的工具,货币加权收益率有一个严重的缺点。一般,客户决定了给予投资经理资金的时间和数量。正如我们所看到的,那些决策可能会显著影响投资经理货币加权收益率的结果。然而,我们评价的一个基本原则是我们只能根据个体自身的行为或者是受其所掌控的行为来对于一个人或一个实体进行评估(判断)。一个(适当的,好的)评估工具应该将投资经理行为的影响排除出去。下面一节所介绍的工具,就在这方面十分有效。

2.3.2 时间加权收益率

对于投资过程中资金投入与提取的影响不敏感的投资度量指标是时间加权收益率。在投资经理行业,时间加权收益率是一个优先考虑的业绩度量指标。时间加权收益率(time-weighted rate of return)度量了投资组合中初始的 $\$1$ 投资在所确定的考察期间内的复利增长率情况。与货币加权收益率不同,时间加权收益率不受到投资组合中现金提取或增加投入的影响。“时间加权”这个术语指的是收益率在时间上进行加权平均。为了计算出确切的投资组合的时间加权收益率,我们采取如下三个步骤:

1. 在任何显著的增加投入或提取资金之前立即给出资产组合的价格。根据现金流入和流出的发生时间,将总的评估期限分成几个子期间。
2. 计算每个子期间中投资组合的持有期回报率。
3. 将持有期回报率相互联系或者其进行复利以得到该年的年收益率(该年的时间加权收益率)。如果投资不止一年,对收益率采用几何平均的方法得到在度量期间上的时间加权收益率。

让我们回到之前货币加权的例子中,并来计算该投资者投资组合的时间加权收益率。在那个

① 在这个特殊的例子中,我们可以通过利用代数的标准解法来求解二次方程 $480x^2 - 220x - 200 = 0$ (其中 $x = 1/(1+r)$)来求解 r 。而一般地,我们将依赖于计算器或者是电子表格软件来计算货币加权收益率。

② 注意计算器或者电子表格所给出的内部收益率是一个期间利率。如果该期间不是年,我们要将该期间利率进行年化。

例子中,我们已经计算了投资组合的持有期回报率,这是上述求解时间加权收益率步骤的第二步。已知投资组合在第一年中获得了15%的收益,在第2年中获得了6.67%的收益,那么在整个两年的评估期中该投资组合的时间加权收益率是多少呢?

我们通过对两个持有期回报率进行几何加权平均来求得时间加权收益率,这就是上述步骤的第三步。几何平均值的计算准确反映了复利增长率的计算。这里,我们通过将 $(1 + \text{每一期的持有期回报率})$ 进行连乘,从而得到在 $t=0$ 时刻投资的\$1在 $t=2$ 时刻的终值。接着,我们对于该乘积开根号再减去1,得到几何平均值。我们将该结果解释为在 $t=0$ 时刻投资组合中的\$1的年度复利增长率。因此,我们有

$$(1 + \text{时间加权收益率})^2 = (1.15)(1.0667)$$

$$\text{时间加权收益率} = \sqrt{(1.15)(1.0667)} - 1 = 10.76\%$$

该资产组合的时间加权收益率是10.76%,与之前的货币加权收益率9.39%相比,它将更多的权重放在了第2年的收益上。由此,我们可以看到为什么投资经理觉得时间加权收益率更有意义。如果一个客户在市场环境不好的时候给予一位投资经理更多的资金,那么该经理的货币加权收益率很可能是非常令人沮丧的。(相反)如果一个客户在市场环境很好的时候增加其投资资金,那么该经理的货币加权收益率很可能是非常鼓舞人心的。而时间加权收益率却排除了这些因素的影响。

在定义计算一个准确的时间加权收益率的步骤中,我们提到投资组合应该在任何明显的增加投入或提取资金之前进行估值。当许多投资组合中现金流活动的数量很大时,这项工作的耗费将是巨大的。于是,我们常常可以通过对某个频率或者是某个固定期间,特别是当增加投入或提取资金与市场变化不相关时,获得一个时间加权收益率的合理估计。估值的频率取得越高,估值的结果也就越精确。通常我们采用的是按日估值。假设一个投资组合在一年中每天进行估值。为了计算其在一年中的时间加权收益率,我们首先计算每天的持有期回报率:

$$r_t = \frac{MVE_t - MVB_t}{MVB_t}$$

式中 MVB_t 等于在 t 时刻这天开始时的市场价格(开盘价),而 MVE_t 为 t 时刻这天结束时的市场价值(收盘价)。我们将所计算的365天的日收益率,记为 $r_1, r_2, \dots, r_{365}$ 。于是,我们通过下面的方式: $(1 + r_1)(1 + r_1) \cdots (1 + r_{365}) - 1$,从每天的持有期回报率得到该年的年收益率。如果对于投资组合的资金提取和增加只发生在每天结束的时刻,那么这个年收益率就是该年精确的时间加权收益率。不然,它仅为时间加权收益率的一个近似估计。

如果我们有一组许多年的数据,可以按上面的方法,分别在每年计算一个时间加权收益率。如果 r_i 表示第 i 年的时间加权收益率,我们可以将 N 年的年收益率几何平均来计算年化的时间加权收益率,具体如下:

$$r_{TW} = [(1 + r_1)(1 + r_2) \cdots (1 + r_N)]^{1/N} - 1$$

例2-4 举例说明了时间加权收益率的计算方法

例2-4 时间加权收益率

斯特鲁贝克(Strubeck)公司发起了一个针对其职员的养老金计划。该计划中的一部分是公司内部所管理的股票投资组合,剩下的另一部分是委托给超级信托(Super Trust)公司管理的。作为斯特鲁贝克公司的投资主管,你想要分别评估一下过去四个季度公司内部和超级信托公

司管理的投资组合的业绩情况。你已经将投资组合所有的现金流出和流入整理为每季季初的状态。表 2-5 汇总了所有现金流的流入与流出情况以及对于两个投资组合的估值。在表中，期末值是在季初的现金流入和流出发生之前的投资组合的价值。投资额为该投资组合经理所负责投资的金额。

表 2-5 斯特鲁贝克 (Strubeck) 组织内部管理账户和超级信托公司账户中的现金流

	季度			
	1	2	3	4
公司内部账户				
期初值	\$4 000 000	\$6 000 000	\$5 775 000	\$6 720 000
期初现金流入 (流出)	\$1 000 000	(\$500 000)	\$225 000	(\$600 000)
投资额	\$5 000 000	\$5 500 000	\$6 000 000	\$6 120 000
期末值	\$6 000 000	\$5 775 000	\$6 720 000	\$5 508 000
超级信托公司账户				
期初值	\$10 000 000	\$13 200 000	\$12 240 000	\$5 659 200
期初现金流入 (流出)	\$2 000 000	(\$1 200 000)	(\$700 000)	(\$400 000)
投资额	\$12 000 000	\$12 000 000	\$5 240 000	\$5 259 200
期末值	\$13 200 000	\$12 240 000	\$5 659 200	\$5 469 568

基于上述信息，回答下面问题：

- (1) 计算公司内部账户的时间加权收益率。
- (2) 计算超级信托公司账户的时间加权收益率。

解：

(1) 为了计算公司内部账户的时间加权收益率，我们先计算账户的季度持有期回报率，然后将其联系起来转化成年利率。公司内部账户的时间加权收益率为 27%，具体计算如下：

一季度持有期回报率： $r_1 = (\$6\,000\,000 - \$5\,000\,000) / \$5\,000\,000 = 0.20$
二季度持有期回报率： $r_2 = (\$5\,775\,000 - \$5\,500\,000) / \$5\,500\,000 = 0.05$
三季度持有期回报率： $r_3 = (\$6\,720\,000 - \$6\,000\,000) / \$6\,000\,000 = 0.12$
四季度持有期回报率： $r_4 = (\$5\,508\,000 - \$6\,120\,000) / \$6\,120\,000 = -0.10$

$(1 + r_1)(1 + r_2)(1 + r_3)(1 + r_4) - 1 = (1.20)(1.05)(1.12)(0.90) - 1 = 0.27$ 或 27%

(2) 由超级信托公司管理的账户其时间加权收益率为 26%，计算如下：

一季度持有期回报率： $r_1 = (\$13\,200\,000 - \$12\,000\,000) / \$12\,000\,000 = 0.10$
二季度持有期回报率： $r_2 = (\$12\,240\,000 - \$12\,000\,000) / \$12\,000\,000 = 0.02$
三季度持有期回报率： $r_3 = (\$5\,659\,200 - \$5\,240\,000) / \$5\,240\,000 = 0.08$
四季度持有期回报率： $r_4 = (\$5\,469\,568 - \$5\,259\,200) / \$5\,259\,200 = 0.04$

$(1 + r_1)(1 + r_2)(1 + r_3)(1 + r_4) - 1 = (1.10)(1.02)(1.08)(1.04) - 1 = 0.26$ 或 26%

公司内部投资组合的时间加权收益率要高于超级信托公司投资组合时间加权收益率 100 个基点。

在做完该练习之后，我们再来看一个更加具体的例子。

例 2-5 同时考虑时间加权收益率和货币加权收益率

你的任务是计算沃布莱特 (Walbright) 基金在 2003 年的投资业绩情况。具体事实如下:

- 在 2003 年 1 月 1 日, 沃布莱特基金的市场价值为 \$100 000 000。
- 在 2003 年 1 月 1 日至 2003 年 4 月 30 日期间, 该基金投资的股票获得了 \$10 000 000 的资本盈利。
- 在 2003 年 5 月 1 日, 该基金投资的股票发放了总额 \$2 000 000 的红利。全部红利被用来再投资于另外的股票。
- 因为该基金表现优异, 机构投资者在 2003 年 5 月 1 日又对沃布莱特基金增加了 \$20 000 000 的额外资金投入, 使得其管理的资产规模达到了 \$132 百万 (\$100 + \$10 + \$2 + \$20)。
- 在 2003 年 12 月 31 日, 沃布莱特收到了总额为 \$2 640 000 的红利。该基金在 2003 年 12 月 31 日的市场价值 (不包括 \$2 640 000 的红利收入) 为 \$140 000 000。
- 该基金在 2003 年期间没有其他的期间现金支付。

基于上述已知信息, 回答下列问题:

- (1) 计算沃布莱特基金的时间加权收益率。
- (2) 计算沃布莱特基金的货币加权收益率。
- (3) 解释时间加权收益率与货币加权收益率之间的差别。

解:

(1) 因为期间的现金流发生在 2003 年 5 月 1 日, 所以我们必须计算两个期间收益率, 然后将它们合并为一个年收益率。表 2-6 列出了在 1 月 1 日、5 月 1 日和 12 月 31 日的相关市场价值以及相关的期间四个月 (1 月 1 日~5 月 1 日) 和八个月 (5 月 1 日~12 月 31 日) 的持有期回报率。

表 2-6 沃布莱特基金的现金流

2003 年 1 月 1 日	初始投资组合的价值 = \$100 000 000
2003 年 5 月 1 日	在增加投资之前的红利收入 = \$2 000 000
	期末投资组合的价值 = \$110 000 000
	持有期回报率 = $(\$2 + \$10) / \$100 = 12\%$
	新增投资 = \$20 000 000
2003 年 12 月 31 日	红利收入 = \$2 640 000
	期末投资组合的价值 = \$140 000 000
	持有期回报率 = $(\$2.64 + \$140 - \$132) / \$132 = 8.06\%$

现在我们必须根据几何平均, 将四个月和八个月的收益率转变成年收益率。我们计算时间加权收益率如下:

时间加权收益率 = $1.12 \times 1.0806 - 1 = 0.2103$

在这个例子中, 我们计算出一个一年 21.03% 的时间加权收益率。四个月和八个月的期间合并起来等于一年。(仅当 1.12 和 1.0806 分别代表一整年的收益率时, 对于乘积 1.12×1.0806 开平方根才是恰当的。)

(2) 为了计算货币加权收益率, 我们需要求解使得现金流出 (购买) 的现值与现金流出 (红利和未来收益) 的现值相等时的贴现率。该基金的初始市场价值和所有的新增资金投入被视为现金流出。(将它们认为是支出。) 提款、收入和期末基金的市场价值被计为现金流入。(期末市场价值是投资者在对该基金进行清算时所获得的收入额。) 因为期间现金流发生在四个月的期间。我们必须求解四个月的内部收益率。表 2-6 中列出了现金流序列及其发生时间。

现值公式（以百万美元计）如下：

$$\begin{aligned} PV(\text{现金流出}) &= PV(\text{现金流入}) \\ \$100 + \frac{\$2}{(1+r)^1} + \frac{\$20}{(1+r)^1} &= \frac{\$2}{(1+r)^1} + \frac{\$2.64}{(1+r)^3} + \frac{\$140}{(1+r)^3} \end{aligned}$$

等式的左侧表示该基金的投资或者现金流出：一笔 \$100 000 000 的初始投资加上 \$2 000 000 的再投资红利以及额外的 \$20 000 000 的新增投资（同时发生在第一个四个月期间的期末，因此，这两项分母的指数为 1）。等式的右侧表示收入或者现金流入：在第一个四个月期间的 \$2 000 000 红利加上 \$2 640 000 的红利以及最终 \$140 000 000 的市场价值（同时发生在第三个四个月期间的期末，因此，这两项分母的指数为 3）。在第二个四个月期间没有现金流发生。我们可以将所有的项都移到等号的右边，并按照时间的先后重新排序。经过化简后，我们有：

$$0 = -\$100 - \frac{\$20}{(1+r)^1} + \frac{\$142.64}{(1+r)^3}$$

利用电子表格或者能够计算内部收益率的计算器，我们分别输入 $t=0$ ， $t=1$ ， $t=2$ 和 $t=3$ 时刻的净现金流 - \$100、- \$20、0 和 \$142.64。[⊖]通过这两个工具任一个的计算，我们得到一个四个月的内部收益率 6.28%。最便捷的将其年化的方法是将结果乘以 3。而更精确的方法是 $(1.0628)^3 - 1 = 0.20$ 或者 20%。

(3) 在这个例子中，时间加权收益率（21.03%）比货币加权收益率（20%）要大。沃布莱特基金在八个月中的表现（即当基金拥有更多股份时）相对于其整体表现较差。这一事实表现在其货币加权收益率相比时间加权收益率要更小，因为货币加权收益率对于投资组合资金的提取和新增投入的时间和数额是非常敏感的。

准确地度量投资组合的收益率在评估投资组合经理业绩的过程中是非常重要的。然而在考虑收益的同时，分析师们还必须权衡风险。当我们做完例 2-4 时，我们没有因为公司内部管理的时间加权收益率更高，就草率地得出公司内部管理优于超级信托公司的结论。如果我们关注风险，我们可以讨论经过风险调整后的业绩，并与之前的结果加以比较（只是更加谨慎）。在之后的几章中，我们将会讨论夏普比率——一个重要的经过风险调整后的业绩度量指标，我们会将其应用于投资经理的时间加权收益率中。至今，我们已经介绍了度量一个投资组合收益率的一些重要工具。

2.4 货币市场收益率

在我们讨论内部收益率和净现值的过程中，我们曾把资本的机会成本视为一个市场确定的利率。在这一节中，我们将开始通过考虑短期债券市场来对现实市场中的贴现现金流分析进行讨论。

为了理解债券市场中不同方式表述的收益率，我们必须讨论货币市场工具收益率报价的传统。货币市场（money market）是指由短期债务工具（一年或小于一年到期）所组成的市场。一些工具要求发行者偿还出借人的本金和利息，而另外一些所谓的纯贴现工具（pure discount instruments）只要求支付利息，该利息等于所需偿还额和所借本金之差。

⊖ 根据传统，我们在现金流出前加上符号，并且我们在 $t=2$ 处需要 0 作为一个占位符。

在美国货币市场中,纯贴现工具的经典例子为由美国联邦政府发行的国库券(Treasury bill, T-bill)。其面值(face value)为美国政府承诺到期偿还给国库券投资者的金额。在购买国库券的过程中,投资者支付其面值与贴现值之差,并在到期日收到其面值。贴现值(discount)是从面值中减去后得到国库券价格的数额。因为投资者会在到期时收到其面值,故贴现值就成为累计的利息。因此,如果投资者持有国库券到期,他将获得与贴现值相等的一笔美元收益。至今,国库券仍是美国货币市场中重要的货币市场工具(种类)之一。其他的货币市场工具种类包括商业票据和银行承兑(这些是贴现工具)以及可转让定期存单(这是付息工具)。对于这些工具中每一种所组成的市场都有其自身的价格或者收益率报价传统。本节的余下部分将对与国库券和其他货币市场工具的报价传统进行一些介绍。在大部分的例子中,收益率的报价必须在进行调整后才能应用于其他的现值问题之中。

纯贴现工具(如国库券)的报价与美国政府债券的报价是不同的。国库券的报价是基于银行贴现基准(bank discount basis),而不是价格基准。银行贴现基准是指一种报价传统,它基于一年360天的标准将贴现值年化为债券面值的一定百分比。基于银行贴现基准的收益率被计算如下:

$$r_{BD} = \frac{D}{F} \frac{360}{t} \quad (2-3)$$

式中 r_{BD} ——基于银行贴现基准的年化收益率;

D ——美元贴现值,它等于国库券面值 F 与其购买价格 P_0 的差价;

F ——国库券的面值;

t ——距离到期日实际剩余的天数;

360——银行对于一年总天数的传统规定。

银行贴现收益率(经常被简单地称为贴现收益率)将美元贴现值 D 从面值中取出来,然后表达为国库券面值(不是其价格)的一个分数。然后这个分数被乘以一年中长度为 t 的期间数量(为 $360/t$),其中一年被假定为有360天。以该方式进行年化,假定了以单利计算(不是以复利计算)。看下面的一个例子。

例 2-6 银行贴现收益率

假设一份国库券具有面值(或者票面价值) \$100 000, 其在150天后到期,目前售价 \$98 000。那么,它的银行贴现收益率是多少?

解:

对于这个例子,其美元贴现值 D 为 \$2 000。基于银行贴现基准的收益率为 4.8%, 利用式(2-3)计算如下:

$$r_{BD} = \frac{\$2\,000}{\$100\,000} \frac{360}{150} = 4.8\%$$

银行贴现公式将国库券美元贴现值从面值或票面价值分离出来,并表示为其的一个分数,2%,然后利用因子 $360/150 = 2.4$ 将其年化。贴现工具如国库券的价格是用贴现收益率报价的,因此我们经常要将贴现收益率转化为价格。

假设我们知道银行贴现收益率为 4.8%,但是不知道其价格。那么我们可以求解其美元贴现值 D , 如下:

$$D = r_{BD} F \frac{t}{360}$$

已知 $r_{BD} = 4.8\%$ ，于是美元贴现值为 $D = 0.048 \times \$100\,000 \times 150/360 = \$2\,000$ 。一旦我们计算出了其美元贴现值，该国库券的买价即为其面值减去其美元贴现值， $F - D = \$100\,000 - \$2\,000 = \$98\,000$ 。

基于银行贴现基准的收益率并不是一个有意义的度量投资者收益率的指标，这主要有三点原因。首先，该收益率是基于债券面值的，而不是基于其购买价格。投资的收益率应该相对于其投资额进行评估。其次，该收益率进行年化时是基于一年 360 天的标准而非一年 365 的现实进行计算的。最后，银行贴现收益率以单利计息方式进行年化，这忽略了获取的利息再产生利息的机会（复利计息）。

我们可以将例 2-6 进行扩展，来讨论三个经常使用的其他收益率度量方式。第一个是货币工具剩余生存期（在例 2-6 中的国库券中为 150 天）的持有期回报率。该指标确定了一个投资者持有该工具到期时所能获得的收益率；正如在这里所使用的，该度量指标是一个没有年化的收益率（或者是一个期间的收益率）。在固定收益市场中，这一持有期回报率（holding period return）也被称为持有期收益率（holding period yield, HPY）。^①对于在生存期中只有一次现金支付的货币工具来说，其持有期收益率可以表示为：

$$\text{HPY} = \frac{P_1 - P_0 + D_1}{P_0} \quad (2-4)$$

式中 P_0 ——初始购买该工具的价格

P_1 ——该工具到期时所卖出的价格

D_1 ——该工具在到期日所支付的现金分配（即利息）

当我们使用这一表达式来对于付息工具（例如付息债券）计算持有期收益率时，我们需要注意一个重要的问题：购买价格和卖出价格必须包括任何加入到交易价格的应计利息（accrued interest），因为该债券是在利息支付日之间被交易的。应计利息是指卖出者应在最后一个付息日所获的息票利息，但其不是以息票形式收到的，因为下一次付息日发生在卖出交易日之后。^②

对于纯贴现证券，其所有收益由超过其购买价格的支付产生。因为国库券是一个纯贴现工具，它不进行利息支付，因此 $D_1 = 0$ 。因此，持有期收益率是其美元贴现值除以其购买价格， $\text{HPY} = D/P_0$ ，其中 $D = P_1 - P_0$ 。持有期收益率为其他度量指标年化后的数值。对于例 2-6 中的国库券来说，投资 $\$98\,000$ 将会在 150 天后收到 $\$100\,000$ 的支付。利用式（2-4）所计算的该投资的持有期收益率为： $(\$100\,000 - \$98\,000)/\$98\,000 = \$2\,000/\$98\,000 = 2.0408\%$ 。对于这个例子，期间收益 2.0408% 是与 150 天的期间相联系的。如果我们把国库券的收益率作为投资的机会成本，我们就可以利用该 150 天国库券的贴现率 2.0408% 来计算任何其他在 150 天后收到现金流的现值。只有其他现金流和该国库券的风险特征类似时，该方法才是恰当的。如果其他现金流比该国库券的风险更大时，我们可以利用国库券收益率作为一个基准利率，在此基础上我们再加上风险溢价。无论货币的面额是多少，该持有期收益率的公式都是相同的。

第二个度量收益率的指标是有效年收益率（effective annual yield, EAY）。有效年收益率将持

① 当讨论总收益时（包括价格变动和利息收入的回报），债券市场的参与者经常使用术语“收益率”表示到期收益率。而在其他的情况下，收益率只表示利息收入的回报（如在当前收益率中，就是指年利息除以价格）正如本书和许多其他作者所使用的，持有期收益率是债券市场中对于持有期回报率、总回报和期间回报的同义词。

② 含有应计利息的价格被称为全价（full price）。交易价格的报价“干净”（是指不含有应计利息），但是一般假设，应计利息是被加入到买价之中的。想要了解更多关于应计利息的内容，参见 Fabozzi（2004）。

有期收益率加上1，然后将其复利至一年，再减去1使之变为年化的收益率。这样做能使该指标反映由利息再投资所获得的利息。[Ⓐ]

$$EAY = (1 + HPY)^{365/t} - 1 \tag{2-5}$$

在我们的这个例子中，我们可以求解有效年收益率如下：

$$EAY = (1.020\,408)^{365/150} - 1 = 1.050\,388 - 1 = 5.038\,8\%$$

该例子揭示了一个普遍的规律：银行贴现收益率要比有效年收益率小。

第三种度量收益率的指标是货币市场收益率（money market yield）（也被称为大额可转让存单等价收益率（CD equivalent yield））。该方法使得国库券收益率的报价能够与以360天为计息基准的付息货币市场工具的收益率报价进行比较。一般地，货币市场收益率等于年化的持有期收益率；假定一年360天， $r_{MM} = (HPY)(360/t)$ 。与银行贴现收益率相比，货币市场收益率是根据购买价格来计算的，所以 $r_{MM} = (r_{BD})(F/P_0)$ 。这一等式显示货币市场收益率要比银行贴现收益率要高。在实际应用中，下面的表达式更加有用，因为它不需要知道国库券的价格：

$$r_{MM} = \frac{360r_{BD}}{360 - (t)(r_{BD})} \tag{2-6}$$

对于国库券的例子来说，其货币市场收益率为 $r_{MM} = (360)(0.048)/[360 - (150)(0.048)] = 4.898\%$ 。[Ⓑ]

表2-7总结了我們至今所讨论的三种收益率的度量指标。

表2-7 三个常用的收益率度量指标

持有期收益率（HPY）	有效年收益率（EAY）	货币市场收益率（大额可转让存单等价收益率）
$HPY = \frac{P_1 - P_0 + D_1}{P_0}$	$EAY = (1 + HPY)^{365/t} - 1$	$r_{MM} = \frac{360r_{BD}}{360 - (t)(r_{BD})}$

下面这个例子将帮助你巩固对于这几个收益率度量指标概念的学习。

例2-7 使用适当的贴现率

你需要求解在150天之后收到的一笔\$1 000现金流的现值。你决定参考150天到期国库券所确定的相关利率来计算该现值。你已经求得一系列150天国库券的收益率。表2-8列出了这些信息。

哪一个或哪几个收益率适合用来求解在150天后收到的\$1 000的现值？

解：

持有期收益率是合适的，并且我们也可以利用经过将货币市场收益率和有效年收益率转换后的持有期收益率。

持有期收益率（2.040 8%）。该收益率确实是我们所需要的。因为它对应于150天的期间，我们可以直接运用它来求解在150天后收到的\$1 000的现值。（记得之前的原则规定，贴现率必须与其对应的时间期间相一致。）故该现值为：

表2-8 短期货币市场收益率

持有期收益率	2.040 8%
银行贴现收益率	4.8%
货币市场收益率	4.898%
有效年收益率	5.038 8%

Ⓐ 有效年收益率在货币时间价值一章中被称为有效年利率（等式1-5）。
Ⓑ 一些国家的市场使用货币市场收益率公式，而不是银行贴现收益率公式，来对诸如国库券这类贴现工具的收益率进行报价。在加拿大，是基于一年365天的货币市场公式对于国库券收益率进行报价。而德国到期日小于一年的贴现票据的收益率和法国国库券（BTF，T-bill）是基于一年360天的货币市场公式来计算的。

$$PV = \frac{\$1\,000}{1.020\,408} = \$980.00$$

现在我们可以来看一下为什么其他收益率度量指标不适合或者不易于应用。

- 银行贴现收益率 (4.8%)。我们不应该使用该收益率度量指标来计算现金流的现值。正如之前所述, 银行贴现收益率是基于国库券的面值计算得来的, 而不是国库券的价格。
- 货币市场收益率 (4.898%)。要使用货币市场收益率, 我们需要将其除以 (360/150) 转化为 150 天的持有期收益率。在获得持有期收益率 $0.048\,98 / (360/150) = 0.020\,408$ 后, 我们使用该值来像前面一样对 \$1 000 进行贴现。
- 有效年收益率 (5.038 8%)。由于该收益率已经年化了, 所以它必须经过调整, 来与现金流所发生的时点相匹配。我们可以从有效年收益率中获得持有期收益率, 具体计算如下:

$$(1.050\,388)^{150/360} - 1 = 0.020\,408$$

回忆一下, 当我们求得有效年收益率时, 其指数为 365/150, 或者是一年 365 天中 150 天期间的个数。为了将有效年收益率缩小到 150 天的收益率, 我们使用之前用来年化时指数的倒数来对于其进行计算。

在例 2-7 中, 我们将两个短期年收益率的度量指标转化为 150 天的持有期收益率。这是一类的转化方式。我们也经常将期间利率转化为年利率。这类问题同样也在货币市场和长期债券市场中存在。作为一个例子, 许多债券 (长期债务工具) 都是每半年支付利息的。而债券投资者所计算的债券内部收益率被称为到期收益率 (YTM)。如果半年到期收益率为 4%, 我们怎样将其年化? 一个准确的方法为, 通过考虑复利, 将其计算为 $(1.04)^2 - 1 = 0.081\,6$ 或者为 8.16%。这个值就是我们之前所称的有效年收益率。然而, 在美国债券市场中所使用的一个方法为, 将半年期的到期收益率翻倍: $4\% \times 2 = 8\%$ 。根据该方法所计算得到的到期收益率, 没有考虑复利, 称为债券等价收益率 (bond-equivalent yield)。而将半年期收益率翻倍所获得的收益率是基于债券等价基准 (bond-equivalent basis) 而得到的。在实务中, 8% 的这个结果常会被简称为债券的到期收益率。在货币市场中, 如果我们为了使得所计算的结果可以与债券到期收益率相比较, 而将六个月期间的收益率翻倍来将其年化, 我们也可以将该计算结果称为债券等价收益率。



统计学概念和市场收益率

3.1 引言

统计方法为我们提供了强大的分析数据，以及从中获得推断结论的一组工具。无论我们分析的是资产收益率、盈利增长率、商品价格还是其他金融数据，统计工具都能帮助我们量化并表现出这些数据的重要特征。在本章中，我们将介绍一些描述与分析数据的基本方法。这些方法在统计学的分支中被称为描述性统计。本章将提供一组在不同投资环境中应用的实用统计学概念和工具。正如本章的章名所反映的，我们要介绍的一个主题就是描述收益率分布的统计方法。[⊖]我们将探讨收益率分布的 4 个性质：

- 收益率集中于何处 [集中趋势 (central tendency)]
- 收益率距其中心位置的偏离有多远 [离散度 (dispersion)]
- 收益率的分布是否是对称的还是偏向一边的 [偏度 (skewness)]
- 收益率是否容易出现极端结果的情况 [峰度 (kurtosis)]

这些概念一般也会同样地被应用到其他类型数据的分布之中。

本章的内容组织如下。在第 2 节给出了一些基本概念的定义之后，在第 3 节和第 4 节中我们会对数据的表述进行讨论。第 3 节讨论的是如何用表格形式来表现数据，第 4 节讨论的是如何用图像来表现数据。接着，我们转而讨论如何对数据分布的情况进行定量描述：第 5 节重点讨论度量数据集中到哪里的量化指标，或数据中心趋势的度量。第 6 节介绍其他描述数据位置的度量指标。第 7 节介绍度量数据离散度的量化指标。第 8 节和第 9 节讨论其他一些能够提供数据更加精确（详细）描述的度量指标。第 10 节讨论如何将第 5 节中介绍的概念应用于实际投资中。

3.2 一些基本概念

在开始本章有关统计学的学习之前，先纵览一下该领域的全貌可能会对我们的学习有所帮

⊖ 伊博森协会 (Ibbotson Associates) (www.ibbotson.com) 慷慨地提供了本章中所使用的大部分数据。我们同样也采用了 Dimson, Marsh 和 Staunton (2002) 的历史数据和他们对于世界金融市场的研究以及其他的一些数据来源。

助。下面，我们将简单地介绍一下统计学的研究范围及其分支学科。我们会解释总体和样本的概念。另外，我们将会看到数据的不同类型对其度量方式及其在分析时所采用的适当的统计方法产生影响。最后，我们将讨论基本的数据度量类型。

3.2.1 统计学的本质

统计 (statistics) 这一术语可以有两种宽泛的含义，一种是指数据，另一种是指方法。一家公司最近 20 个季度的平均每股盈利 (earnings per share, EPS) 或者最近 10 年的平均收益率就是统计。我们也可以通过分析历史的 EPS 来预测未来的 EPS，或者利用公司过去的收益率来推断其所面临的风险。我们采用的所有这些收集与分析数据的方法也被称为统计。

统计方法包括描述性统计和统计推断 (推断性统计)。描述性统计 (descriptive statistics) 研究的是如何有效地概括数据以反映包含大量数据的集合的重要特征。描述性统计通过对大量的数据细节进行整合，将数据转化为有用的信息。统计推断 (statistical inference) 所涉及的是如何通过实际观测到的小的数据集合来对大的数据集合的情况进行预测、估计和判断。统计推断的基础是概率论。统计推断和概率论将在后面的几章中进行讨论。本章的重点只放在描述性统计的讨论上。

3.2.2 总体和样本

在整个统计学的学习中，我们要明确区分总体和样本这两个概念之间的区别。在这一节中，我们将在解释这两个术语的同时，介绍与其相关的“参数”和“样本统计量”这两个术语。[⊖]

- **总体的定义。**一个总体 (population) 是指一组特定数据中的所有元素。

任何对于总体特征的描述性度量指标均被称为参数 (parameter)。虽然对总体进行刻画参数可能会有很多，但是投资分析师们一般只关注其中的一部分，如投资收益率的均值、极差以及方差。

即使我们能够观测总体中的所有元素，这样做需要花费的高昂的时间和资金成本也会令我们望而却步。例如，如果总体是全球无线通信的全体用户，而分析师想要了解他们的购买计划，那么分析师将会发现，如果要观测整个总体的情况，花费将是巨大的。这时，分析师可以通过从总体中取出一个样本来对该问题进行分析。

- **样本的定义。**一个样本 (sample) 就是总体的一个子集。

在抽取一个样本时，分析师往往希望该样本能够反映总体的特征。在统计学领域中，关于采用适当的方法来使得抽出的样本能够达到较好地反映总体目的的研究被称为抽样 (调查)。在本章的后面部分，我们会介绍一些具体的抽样方法。

之前，我们在提到数据的时候已经涉及了统计学的相关内容。正如参数是描述总体特征的度量指标，样本统计量 (简称为统计量) 是描述样本特征的度量指标。

- **样本统计量的定义。**一个样本统计量 (或统计量, sample statistic or statistic) 是由样本计算得来的，用来描述样本特征的一个量。

⊖ 本章会介绍许多统计学的概念和公式。为了使得它们能更易于识别，我们将一些重要的内容用项目符号列举出来。

我们在本章的大部分内容都会在这个定义下解释和介绍统计量的用途。这个概念同样在统计推断中也是非常重要的，因为它涉及如何应用样本统计量来对于未知总体参数进行估计的问题。

3.2.3 度量尺度

为了选择一个恰当的统计方法来描述和分析数据，我们需要区分不同的度量尺度（measurement scales）或者说是测量标准。所有数据都会用下面四种度量尺度之一进行度量：名义型、顺序型、区间型或者比率型。

名义尺度（nominal scales）代表最简单的度量标准：它对数据进行分类但不进行排序。如果我们根据共同基金不同的投资策略用整数对其进行编号，那么数字1可能代表小市值价值型基金，数字2代表大市值价值型基金，诸如此类，对于每种类型的基金我们都会用一个数字来进行表示。这种名义尺度将基金根据各自的类型进行分类，但并没有进行排序。

顺序尺度（ordinal scales）代表稍强一点的度量标准。顺序尺度将数据分成根据某种特征排序的不同类别。举个例子来说，晨星和标准普尔为共同基金进行的评级服务用的就是顺序尺度，根据共同基金的相对业绩表现，将最差的一类基金评为一星，然后随着业绩向好的方向发展，依次评为二星、三星、四星和五星。

顺序尺度也可以用数字来区分不同的类别。例如，在对平衡型共同基金5年的累积收益率进行排序时，我们可以将数字1分配给业绩最好的前10%的基金，以此类推，从而数字10代表表现最差的后10%的基金。因为顺序尺度能够揭示出排名1的基金比排名2的基金业绩更好，所以顺序尺度比名义尺度要求更高。但是，该尺度并没有告诉我们排名1的基金和排名2的基金之间业绩的差别与排名3和排名4，或者排名9和排名10的基金业绩之间的差别有什么不同。

间隔尺度（interval scales）不仅提供数据的排序，还保证相邻赋值之间的差别是相同的。因此，这些赋值之间是可以进行有意义的加减运算的。例如，摄氏和华氏的尺度就是这类区间型度量尺度。温度在10℃和11℃之间的差别与40℃和41℃之间的差别是相同的，而且我们可以准确地得到 $12^{\circ}\text{C} = 9^{\circ}\text{C} + 3^{\circ}\text{C}$ 。但是，间隔尺度的零点并不表示我们所度量的东西不存在；它不是一个真正的零点或者自然的零。零摄氏度对应的是水的冰点，而不是指没有温度。由于缺乏真正的零点，所以我们不能通过间隔尺度来计算有意义的比率。

举例来说，50℃虽然比10℃大5倍，但是这并不意味着温度高5倍。同样，问卷中的量常会用间隔尺度来进行度量。如果一个投资者被要求用从1（极端风险厌恶）到7（极端风险喜好）中的一个数来度量其风险厌恶程度，那么回答为1和回答为2的投资者之间的风险厌恶差别，与回答为6和回答为7的投资者之间的风险厌恶差别有时就被假设为相同的。当我们做出上述假设时，该数据就是由间隔尺度进行度量的。

比率尺度（ratio scales）代表要求最高的度量标准。该尺度不仅具有区间度量尺度的所有特征，而且拥有一个真正的零点作为数据的原点。对于用比率尺度进行度量的数据，我们计算出的比率是有意义的，而这些比率间的加减运算同样也是有意义的。因此，我们可以应用所有的统计工具来对以比率尺度度量的数据进行处理。收益率和货币一样，都是以比率尺度来度量的量。如果我们有两倍的货币量，那么我们就有两倍的购买力。注意该度量是有一个自然的零值的（零意味着没有货币）。

例 3-1 识别度量尺度

说明下列各项指标所使用的度量尺度：

- (1) 债券发行的信用评级。^①
- (2) 每股现金红利。
- (3) 对冲基金分类类别。^②
- (4) 债券到期年限。

解：

(1) 信用评级是以顺序尺度来度量的。一个评级将发行的债券归到一个类别之中，并且该类别是根据其期望违约概率来进行排序的。但是 AA- 和 A+ 之间期望违约概率的差别并不必然与 BB- 和 B+ 之间期望违约概率的差别相同。也就是说，用字母来表示的信用评级并不是以间隔尺度来度量的。

(2) 每股现金红利是以比率尺度来度量的。对于该变量，0 表示完全没有红利；这是一个真正的零点。

(3) 对冲基金分类类别是以名义尺度来度量的。每个类别中包含的是以相同投资策略进行投资的对冲基金。然而，与债券的信用评级不同，对冲基金分类规则中并没有涉及对其进行排序。因此，该分类规则并不是以顺序尺度来度量的。

(4) 债券到期年限是以比率尺度来度量的。

3.3 用频数分布汇总数据

在这一节中，我们讨论一种最简单的汇总数据的方法——频数分布。

- **频数分布的定义。**一个频数分布 (frequency distribution) 是指一种表格列示数据的方法，它用较少的区间对数据总体进行概括。

频数分布为我们对大量统计数据的分析提供了帮助，并且其能够应用于任何尺度度量的数据之上。

收益率是分析师和投资组合经理用来进行投资决策的最基本要素，我们可以使用频数分布来对其进行汇总与描述。当我们分析收益率时，我们的出发点是持有期收益（也被称为总收益）。

- **持有期收益公式 (holding period return formula)。**在时刻 t (从 $t-1$ 到 t) 的持有期收益 R_t 为：

$$R_t = \frac{P_t - P_{t-1} + D_t}{P_{t-1}} \quad (3-1)$$

① 债券发行的信用评级度量了债券发行者履行债券所承诺支付的本金和利息的能力。例如，评级机构标准普尔，给予每一个发行的债券分配如下的某一个级别，以信用质量的降序排列的具体级别分别为（或者以违约概率的升序排列）：AAA，AA+，AA，AA-，A+，A，A-，BBB，BBB-，BB+，BB，BB-，B，CCC+，CCC-，CC，C，CI，D。想要了解更多关于信用风险的信息，参见 Fabozzi (2004a)。

② “对冲基金”指的是一种具有合法组织结构的投资机构，它所受的监管要小于其他的集资投资机构，如共同基金。根据它们所采用的投资策略的类型可将对冲基金划分成不同的类别。

式中, P_t 表示时刻 t 的每股价格; P_{t-1} 表示时刻 $t-1$ 的每股价格, 该时刻紧接在时刻 t 之前; D_t 表示在时刻 $t-1$ 到时刻 t 之间收到的现金分配。

因此, 时刻 t 的持有期收益为资本利得 (或者损失) 加上现金分配再除以期初的价格。(对于普通股来说, 现金分配为红利; 对于债券来说, 现金分配为息票支付。) 我们可以简单地通过对于时间指标 t 的间隔单位的不同解释, 来利用式 (3-1) 对任何资产的日、周、月或者年的持有期收益分别进行定义。

如式 (3-1) 所定义的持有期收益具有两个重要的特征。首先, 它是一个与时间有关的量。例如, 如果我们观测到的价格之间的时间间隔是以月为单位的话, 那么收益率就是一个月度的数字。其次, 收益率是没有货币单位的。例如, 如果货币是用欧元来计量的话, 那么式 (3-1) 的分子与分母都将以欧元来表示, 结果该比率将不含有任何单位, 因为分子和分母的单位可以相互约去。无论价格是以何种货币来计量的, 这一结果都将成立。^①

在注意到上述事项之后, 我们现在来讨论标准普尔 500 指数持有期收益的频数分布。^② 首先, 我们考察年度收益率; 然后我们再看月度收益率。通过式 (3-1), 我们可以计算出标准普尔 500 在 1926 年 1 月至 2002 年 12 月的年度收益率, 共 77 个年度观测值。月度数据包括 1926 年 1 月至 2002 年 12 月, 共 924 个月度观测值。

我们可以将建立频数分布的基本步骤叙述如下:

• 建立一个频数分布。

1. 将数据以升序排列。
2. 计算出数据的极差。定义极差 = 最大值 - 最小值。
3. 确定频数分布包含的区间数 k 。
4. 确定区间的宽度 (极差/ k)。
5. 通过不断地在数据最小值上加上区间宽度来确定各个区间的端点, 该过程在到达包含最大值的区间时停止。
6. 计算落入每个区间中观测值的个数。
7. 建立一个列示落入从小到大排列的每个区间中观测值数量的表格。

在第 4 步中, 当对区间宽度进行四舍五入时, 我们采用上舍入 (round up) 而不是下舍出 (round down) (如 1.11 的上舍入为 2, 而下舍出为 1), 以保证最后的那个区间中包含数据的最大值。

正如上述步骤所明确表述的, 一个频数分布将数据归入到一系列区间之中。^③ 一个区间 (interval) 是指一对数值, 在这对数值的范围之内会包含落入其中的观测值。每个观测值只会落入一个区间之中, 并且整个区间应该包含数据中所有观测值。实际落入一个给定区间的观测值数量被称为绝对频数 (absolute frequency), 或者简称为频数。频数分布是反映一组区间以及各个区间相对应的频数度量指标的一个列表。

① 然而我们要注意, 如果在持有期收益中的价格和现金分配不是以同一本国货币来表述的话, 我们一般需要在计算持有期收益之前先将这些变量都转化为本国货币。由于汇率会在持有期间内波动, 所以一个资产的持有期收益以不同货币作为单位进行计算时, 往往会得出不同的结果。

② 我们使用伊博森协会提供的标准普尔 500 在 1926 年 1 月至 2002 年 12 月的总收益序列。

③ 区间有时也被称为组别 (classes)、范围 (ranges) 或者箱体 (bins)。

为了说明操作的基本步骤，假设我们有 12 个观测值，以升序排列为：-4.57，-4.04，-1.64，0.28，1.34，2.35，2.38，4.28，4.42，4.68，7.16 和 11.43。最小的观测值为 -4.57，而最大的观测值为 +11.43，因此，极差为 $+11.43 - (-4.57) = 16$ 。如果我们设定 $k = 4$ ，那么区间宽度为 $16/4 = 4$ 。表 3-1 显示了不断增加区间宽度 4 而得到的区间的端点（步骤 5）。

因此，所得到的区间分别为 $[-4.57, -0.57)$ ， $[-0.57, 3.43)$ ， $[3.43, 7.43)$ 和 $[7.43, 11.43)$ 。[⊖]表 3-2 汇总了步骤 5 到步骤 7 的结果。

在实务中，我们可能想要将上述的基本步骤加以简化。例如，我们可能为了解释方便，想让区间以整数开始与结束。另外，我们也需要对区间数 k 的选择予以解释。我们将在对于标准普尔 500 数据建立频数分布的过程中，对于这些问题进行讨论。

我们首先讨论如何对 1926 ~ 2002 年区间段的标准普尔 500 年收益率建立频数分布。在这个区间中，标准普尔 500 收益率的最小值为 -43.34%（在 1931 年），最大值为 +51.99%（在 1933 年）。因此，该数据的极差近似为 $+54\% - (-43\%) = 97\%$ 。现在的问题是我们用来归类观测值的区间数 k 应该选择多少。虽然在统计学的文献中有一些对于 k 选择的建议，但是设置一个有用的 k 值经常涉及对于数据仔细的观察与审慎的判断。那么我们应该要包含多少具体细节呢？如果我们使用过少的区间，那么我们将汇总得过于粗略，从而失去数据相关的特征。而如果我们使用过多的区间，那么我们可能根本就没有起到汇总的效果。

我们可以通过比较划分不同区间宽度所产生的效果，来选择一个适当的 k 值。大量的空白区间可能意味着我们过分想要将数据进行分类，以至于反映了过多的细节。我们可以从一个相对较小的区间宽度开始着手，来观察它是否使得大部分区间是空白的，以及判断与区间宽度相关的 k 值是否取得过大。如果区间大部分都是空白的或者 k 值非常大，那么我们可以考虑逐步增大区间宽度（减小 k 值），直至产生一个能有效汇总数据分布的频数分布。对于标准普尔 500 年度序列来说，1% 的收益率区间宽度将形成 97 个区间，并且其中很多区间是空白的，因为我们只有 77 个年度观测值。我们需要记住建立频数分布的目的在于汇总数据。假如为了解释简单，我们想要使用以整数而不是分数百分比表述的区间宽度。一个 2% 的区间宽度可能会比 1% 的区间宽度拥有更少的空白区间，并且能更有效地汇总数据。一个 2% 的区间宽度将会有 $97/2 = 48.5$ 个区间，我们可以将其四舍五入成 49 个区间。总区间的宽度将为 $2\% \times 49 = 98\%$ 。我们可以肯定，如果我们从最小的整数 -44% 开始，不断增加 2% 的区间，最终区间的端点将位于 $-44.0\% + 98\% = 54\%$ ，该区间就包含了样本中最大的收益率 53.99%。在通过这种方式建立该频数分布过程中，我们也得到以数值 0% 结束和以 0% 开始的两个区间，这样就使得我们可以计算数据中负的和正的收益率的数量。不需要太多的工作量，我们就找到了汇总数据的有效方法。我们将使用 2% 的区间，从 $-44\% \leq R_t < -42\%$ 开始（在表中给出的是“-44.0% 到 -42.0%”）一直到 $52\% \leq R_t \leq 54\%$

表 3-1 区间的端点

$-4.57 + 4.00 = -0.57$
$-0.57 + 4.00 = 3.43$
$3.43 + 4.00 = 7.43$
$7.43 + 4.00 = 11.43$

表 3-2 频数分布

区间	绝对频数
A $-4.57 \leq \text{观测值} < -0.57$	3
B $-0.57 \leq \text{观测值} < 3.43$	4
C $3.43 \leq \text{观测值} < 7.43$	4
D $7.43 \leq \text{观测值} < 11.43$	1

注：这里的区间之间不会相互重叠，所以每个观测值只能唯一地落入其中的某一区间之中。

⊖ $[-4.57, -0.57)$ 的记法是指 $-4.57 \leq \text{观测值} < -0.57$ 。在文中，方括号表示该端点包含在区间内。

结束。表 3-3 给出了标准普尔 500 年度总收益率的频数分布。

表 3-3 中包含了其他 3 种表示数据的有用方法。一旦建立了频数分布，我们就可以计算相对频数、累积频数（也被称为累积绝对频数）和累积相对频数。

表 3-3 1926 ~ 2002 年标准普尔 500 年度总收益率的频数分布

收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)	收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)
-44.0 ~ -42.0	1	1.30	1	1.30	6.0 ~ 8.0	4	5.19	33	42.86
-42.0 ~ -40.0	0	0.00	1	1.30	8.0 ~ 10.0	1	1.30	34	44.16
-40.0 ~ -38.0	0	0.00	1	1.30	10.0 ~ 12.0	3	3.90	37	48.05
-38.0 ~ -36.0	0	0.00	1	1.30	12.0 ~ 14.0	1	1.30	38	49.35
-36.0 ~ -34.0	1	1.30	2	2.60	14.0 ~ 16.0	1	1.30	39	50.65
-34.0 ~ -32.0	0	0.00	2	2.60	16.0 ~ 18.0	2	2.60	41	53.25
-32.0 ~ -30.0	0	0.00	2	2.60	18.0 ~ 20.0	6	7.79	47	61.04
-30.0 ~ -28.0	0	0.00	2	2.60	20.0 ~ 22.0	3	3.90	50	64.94
-28.0 ~ -26.0	1	1.30	3	3.90	22.0 ~ 24.0	5	6.49	55	71.43
-26.0 ~ -24.0	1	1.30	4	5.19	24.0 ~ 26.0	2	2.60	57	74.03
-24.0 ~ -22.0	1	1.30	5	6.49	26.0 ~ 28.0	1	1.30	58	75.32
-22.0 ~ -20.0	0	0.00	5	6.49	28.0 ~ 30.0	1	1.30	59	76.62
-20.0 ~ -18.0	0	0.00	5	6.49	30.0 ~ 32.0	5	6.49	64	83.12
-18.0 ~ -16.0	0	0.00	5	6.49	32.0 ~ 34.0	4	5.19	68	88.31
-16.0 ~ -14.0	1	1.30	6	7.79	34.0 ~ 36.0	0	0.00	68	88.31
-14.0 ~ -12.0	0	0.00	6	7.79	36.0 ~ 38.0	4	5.19	72	93.51
-12.0 ~ -10.0	4	5.19	10	12.99	38.0 ~ 40.0	0	0.00	72	93.51
-10.0 ~ -8.0	7	9.09	17	22.08	40.0 ~ 42.0	0	0.00	72	93.51
-8.0 ~ -6.0	1	1.30	18	23.38	42.0 ~ 44.0	2	2.60	74	96.10
-6.0 ~ -4.0	1	1.30	19	24.68	44.0 ~ 46.0	0	0.00	74	96.10
-4.0 ~ -2.0	1	1.30	20	25.97	46.0 ~ 48.0	1	1.30	75	97.40
-2.0 ~ 0.0	3	3.90	23	29.87	48.0 ~ 50.0	0	0.00	75	97.40
0.0 ~ 2.0	2	2.60	25	32.47	50.0 ~ 52.0	0	0.00	75	97.40
2.0 ~ 4.0	0	0.00	25	32.47	52.0 ~ 54.0	2	2.60	77	100.00
4.0 ~ 6.0	4	5.19	29	37.66					

注：每个区间的下限取的是不严格不等号（≤），而每个区间的上限取的是严格不等号（<）。
资料来源：由伊博森协会 EnCorr Analyzec 产生的频数分布。

- **相对频数的定义。**相对频数（relative frequency）是指每个区间的绝对频数除以整个样本观测值的数量。

累积相对频数（cumulative relative frequency）将随着区间从第一个到最后一个的移动，对于相应的相对频数进行累积（加总）。该指标告诉我们低于每个区间上限的观测值占总观测值的比例。‘仔细观测表 3-3 中给出的频数分布，我们可以看到，第 1 个收益率区间 -44% 到 -42%，只有一个观测值；其相对频数为 1/77 或 1.30%。因为低于 -42% 的观测值只有一个，所以该区间的累积频数为 1。因此，累积相对频数为 1/77 或 1.30%。第 2 个收益率区间有 0 个观测值；因此，其累积频数为 0 加上 1，而其累积相对频数为 1.30%（从前一个区间得来的累积相对频数）。我们可以通过在前一个累积频数的基础上加总（绝对）频数来求得其他区间的累积频数。累积频

数告诉我们小于每个收益率区间上限的观测值的数量。

正如表 3-3 所显示的, 该样本在各个收益率区间中频数的取值为 0 ~ 7 中的一个数。收益率为 -10% ~ -8% 的区间中最多有 7 个观测值。频数第二多的区间为收益率 18% ~ 20% 的 6 个观测值。从累积频数一行中, 我们可以看到, 负收益率的观测值为 23 个。所以正收益率的观测值一定等于 77 - 23, 即 54 个。我们可以将正的和负的收益率的数目分别表示成总数的一个百分比, 并通过该指标来感受投资于股票市场所存在的内在风险。在 77 年的区间中, 标准普尔 500 具有负的年收益率的时间占总体的 29.9% (即 23/77)。这一结果显示在表 3-3 的第 5 列, 它给出了累积相对频数的结果。

频数分布不仅告诉我们大部分观测值集中在何处, 还告诉我们数据分布是否是对称的、有偏的或者尖峰的。在标准普尔 500 的例子中, 我们可以看到, 超过一半的数据都是正的, 而且大部分的年收益率都超过 10% (54 个正的年收益率中只有 11 个 (大约 20%) 集中在 0 到 10% 之间)。

表 3-3 允许我们对于区间数的选择, 特别是与股票收益率有关的区间数的选择问题, 给出了进一步的重要结论。从表 3-3 的频数分布中, 我们可以看到, 落在 -44% 至 -16% 和 38% 至 54% 之间的结果只有 5 个。股票收益率数据经常会出现非常大或非常小的结果。我们可以通过选择一个较小的 k 值, 使得频数分布尾部的收益率区间减少。但是这么做会使得我们失去关于股票市场极端差或者极端好这两种情况的信息。一位风险管理者可能需要知道最差的可能结果, 因此, 他可能想要关于分布尾端 (极端值) 部分的详细信息。此时, 一个具有较大 k 值的频数分布就会对他有用。而一个投资组合经理或者投资分析师同样也可能对于尾部的详细信息感兴趣; 然而, 如果一个经理或者分析师想要获得关于大部分观测值主要集中在何处的信息, 那么他可能会希望区间的宽度更大一些, 例如 4% (从 -44% 开始, 共 25 个区间)。

标准普尔 500 月度收益率的频数分布看上去与年度收益率有较大不同。从 1926 年 1 月至 2002 年 12 月的月度收益率序列共 924 个观测值。收益率从最小的 -30% 左右到最大的 43% 左右。对于这么大规模的月度数据, 我们必须对其进行汇总以获得对于数据总体分布的一个感受, 所以我们将数据分为 37 个 2% 宽度的等间距收益率区间。通过这样的汇总方式所获得的好处是巨大的。表 3-4 列出了最终获得的频数分布。绝对频数出现在第二列中, 而在其后面的是相对频数一行。相对频数一行中的数据都被近似到小数点后两位。而累积绝对频数和累积相对频数分别出现在第 4 列和第 5 列中。

表 3-4 1926 年 1 月至 2002 年 12 月标准普尔 500 月度总收益率的频数分布

收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)	收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)
-30.0 ~ -28.0	1	0.11	1	0.11	-10.0 ~ -8.0	20	2.16	43	4.65
-28.0 ~ -26.0	0	0.00	1	0.11	-8.0 ~ -6.0	30	3.25	73	7.90
-26.0 ~ -24.0	1	0.11	2	0.22	-6.0 ~ -4.0	54	5.84	127	13.74
-24.0 ~ -22.0	1	0.11	3	0.32	-4.0 ~ -2.0	90	9.74	217	23.48
-22.0 ~ -20.0	2	0.22	5	0.54	-2.0 ~ 0.0	138	14.94	355	38.42
-20.0 ~ -18.0	2	0.22	7	0.76	0.0 ~ 2.0	182	19.70	537	58.12
-18.0 ~ -16.0	2	0.22	9	0.97	2.0 ~ 4.0	153	16.56	690	74.68
-16.0 ~ -14.0	3	0.32	12	1.30	4.0 ~ 6.0	126	13.64	816	88.31
-14.0 ~ -12.0	5	0.54	17	1.84	6.0 ~ 8.0	58	6.28	874	94.59
-12.0 ~ -10.0	6	0.65	23	2.49	8.0 ~ 10.0	21	2.27	895	96.86

(续)

收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)	收益率区间 (%)	频数	相对频数 (%)	累积 频数	累积相对 频数 (%)
10.0 ~ 12.0	14	1.52	909	98.38	28.0 ~ 30.0	0	0.00	921	99.68
12.0 ~ 14.0	6	0.65	915	99.03	30.0 ~ 32.0	0	0.00	921	99.68
14.0 ~ 16.0	2	0.22	917	99.24	32.0 ~ 34.0	0	0.00	921	99.68
16.0 ~ 18.0	3	0.32	920	99.57	34.0 ~ 36.0	0	0.00	921	99.68
18.0 ~ 20.0	0	0.00	920	99.57	36.0 ~ 38.0	0	0.00	921	99.68
20.0 ~ 22.0	0	0.00	920	99.57	38.0 ~ 40.0	2	0.22	923	99.89
22.0 ~ 24.0	0	0.00	920	99.57	40.0 ~ 42.0	0	0.00	923	99.89
24.0 ~ 26.0	1	0.11	921	99.68	42.0 ~ 44.0	1	0.11	924	100.00
26.0 ~ 28.0	0	0.00	921	99.68					

注：每个区间的下限取的是不严格不等号（≤），而每个区间的上限取的是严格不等号（<）。相对频数是绝对频数或者累积频数除以总观测数目。
资料来源：由伊博森协会 EnCorr Analyzec 产生的频数分布。

频数分布的优点都明显地反映在表 3-4 中，该表告诉我们大部分的观测值（599/924 = 65%）都位于 -2% 到 +6% 的 4 个区间之中。我们共有 355 个负的收益率和 569 个正的收益率，几乎 62% 的月度收益率都是正的。从最后一列的累积相对频数来看，-2% 到 0% 区间的累积频数表明 38.42% 的数据低于 0% 的收益率上限，这意味着 38.42% 的观测值低于 0% 的收益率水平。我们同样能够看到，并没有很多观测值大于 12% 或者小于 -12%。我们要注意年度收益率的频数分布和月度收益率的频数分布不是直接可比的。通常，我们应该能预计到以更短的区间进行度量的收益率（例如月收益率）要小于在较长时间区间中进行度量的收益率（如年收益率）。

接着，我们就对于 1900 ~ 2000 年间 16 个主要股票市场的通胀调整后的平均收益率建立频数分布。

例 3-2 建立一个频数分布

不同国家的投资者在长期中的股票收益情况是如何的？为了回答这个问题，我们可以直接对平均年度收益率进行检验。[⊖]然而，名义收益率水平的数值依赖于货币购买力的变化，并且各国还面临着不同的价格通胀情况。因此，更好的方法是我们对不同国家投资者获得的平均实际收益率或者经过通胀调整后的收益率进行比较。蒂姆森（Dimson）、马什（Marsh）和斯汤顿（Staunton）（2002）提供了 1900 ~ 2000 年这 101 年间 16 个国家的资产收益率的权威数据。表 3-5 摘取了一些他们提供的经过通胀调整后的平均收益率的数据。

表 3-5 真实（经过通胀调整之后的）股票收益率：16 个主要股票市场（1900 ~ 2000 年）

国家	算术平均 (%)	国家	算术平均 (%)	国家	算术平均 (%)
澳大利亚	9.0	爱尔兰	7.0	西班牙	5.8
比利时	4.8	意大利	6.8	瑞典	9.9
加拿大	7.7	日本	9.3	瑞士	6.9
丹麦	6.2	荷兰	7.7	英国	7.6
法国	6.3	南非	9.1	美国	8.7
德国	8.8				

资料来源：蒂姆森、马什和斯汤顿（2002）中的表 4-3，瑞士股票市场的数据是从 1911 年开始的。

⊖ 一组数值的平均数或者算术平均数等于这组数值的总和除以该组数值的总个数。例如，为了求解 101 个年收益率的算术平均数，我们将 101 个年收益率进行加总然后除以总数 101。算术平均数的概念将在之后章节的常用统计学概念中进行详细的介绍。

表 3-6 将表 3-5 中的数据汇总成 6 个跨度从 4% ~ 10% 的区间。

表 3-6 平均真实股票收益率的频数分布				
收益率区间 (%)	频数	相对频数 (%)	累积频数	累积相对频数 (%)
4.0 ~ 5.0	1	6.25	1	6.25
5.0 ~ 6.0	1	6.25	2	12.50
6.0 ~ 7.0	4	25.00	6	37.50
7.0 ~ 8.0	4	25.00	10	62.50
8.0 ~ 9.0	2	12.50	12	75.00
9.0 ~ 10	4	25.00	16	100.00

注：相对频数以总和为 100% 的比率表示。

正如表 3-6 所示，各国平均实际股票收益率之间存在着巨大的差别。3/4 的观测值都落在 3 个区间（6.0% ~ 7.0%，7.0% ~ 8.0%，9.0% ~ 10.0%）中的其中一个。大部分的平均实际股票收益率都位于 6.0% ~ 10% 之间；收益率小于 6.0% 的累积相对频数只有 12.50%。

3.4 数据的图形表示

数据的图形表示能使我们直观地发现数据的重要特征。例如，我们可以从图中看出分布是否是对称，并且该结果会对我们用什么概率分布来描述数据产生重大的影响。在这一节中，我们将讨论用图像表示数据的方法，具体包括直方图、频数多边形和累积频数分布图。我们将把表 3-3 或者表 3-4 中标准普尔 500 频数分布所包含的信息用所有这些图形表示出来。

3.4.1 直方图

直方图是与频数分布等价的图像表示方法。

- **直方图的定义。**直方图（histogram）是指用来描述由数据分类而得到的频数分布的条状图。

可视化表示的优点在于：我们可以迅速地看到大部分观测值所在的位置。为了了解直方图是如何制作的，我们来看表 3-3 中的收益率区间 $18\% \leq R_t \leq 20\%$ 。这个区间的绝对频数是 6。因此，我们在图上表示该收益率区域的水平轴上方画出一个条型或者长方形，其高度为 6。对于所有其他收益率区间不断重复该过程，最终就得到了一个直方图。图 3-1 显示了由 1926 ~ 2002 年标准普尔 500 年度总收益率所绘制的直方图。

在图 3-1 的直方图中，每个条状物的高度代表每个收益区间的绝对频数。收益率区间 $-10\% \leq R_t \leq 8\%$ 其频数为 7，它在直方图中以最高的条状物表示。因为区间边界之间没有空隙，所以直方图中的条状物之间也没有间隔。由于许多收益率区间的频数为 0，因此，这些区间在直方图中没有高度。

图 3-2 反映的是由标准普尔 500 月度收益率分布所绘制的直方图。这个直方图比图 3-1 中年收益率的直方图稍稍更加对称一些，也更加接近于钟的形态。

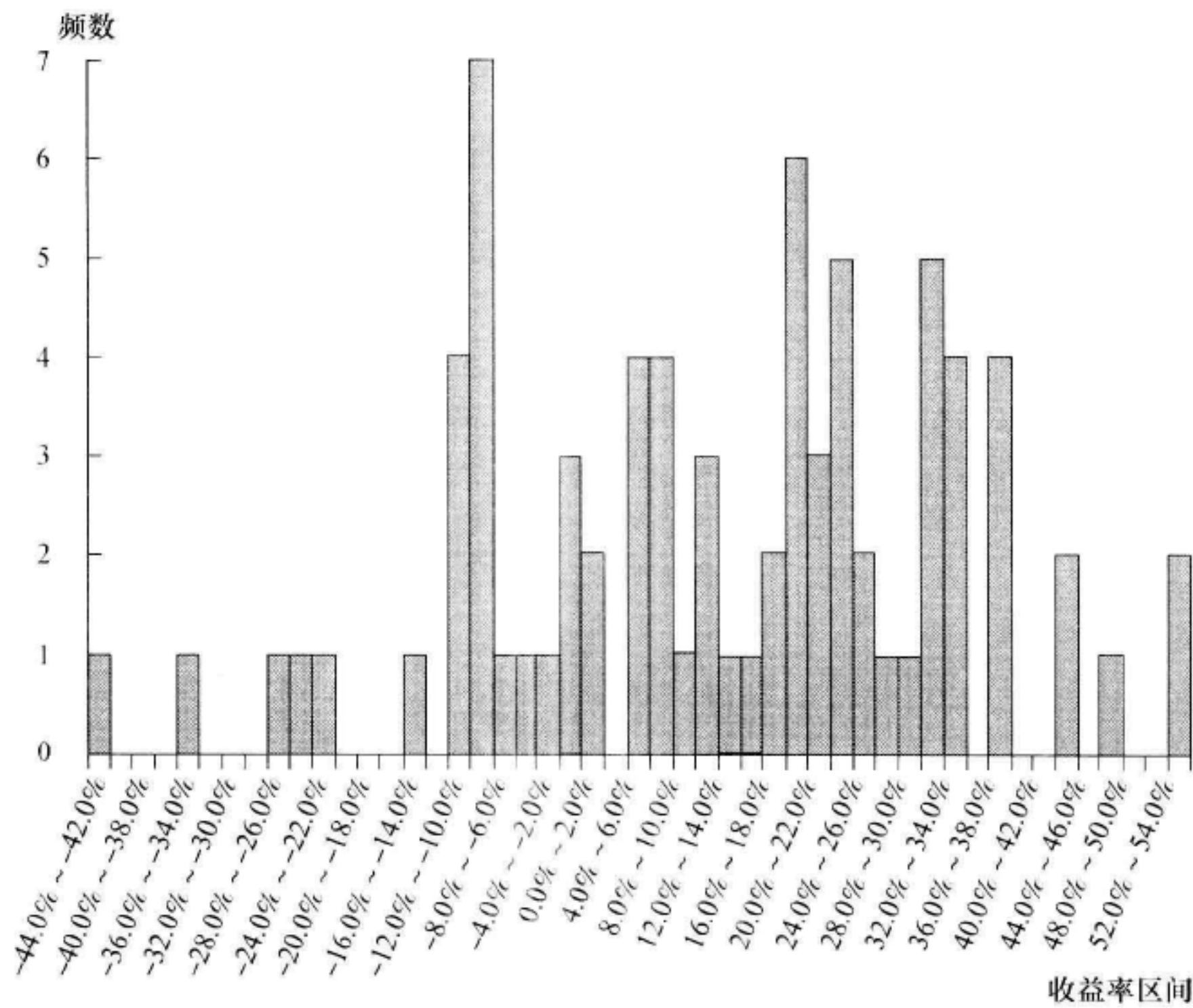


图 3-1 由 1926 ~ 2002 年的标准普尔 500 年度总收益率所绘制的直方图
资料来源：伊博森 EnCorr Analyzer™。

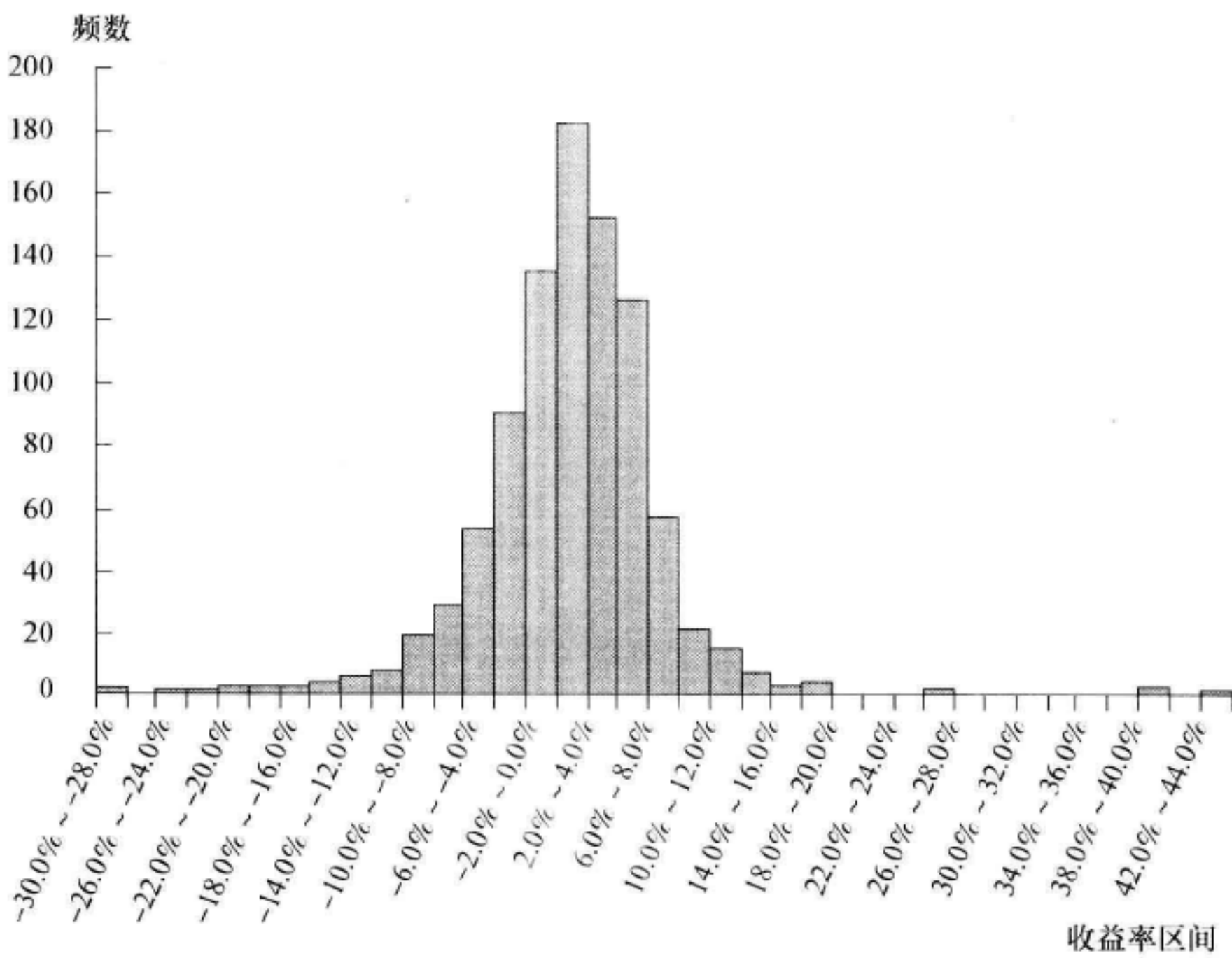


图 3-2 由 1926 年 1 月 ~ 2002 年 12 月标准普尔 500 月度总收益率所绘制的直方图
资料来源：伊博森 EnCorr Analyzer™。

3.4.2 频数多边形和累积频数分布图

另外两种用来描述数据的图形工具是频数多边形和累积频数分布图。为了建立频数多边形 (frequency polygon)，我们标出 x 轴上每个区间的中点，并在 y 轴上画出该区间的绝对频数；然后将相邻的点用直线连接起来。图 3-3 展现的是用 1926 年 1 月至 2002 年 12 月标准普尔 500 924 个月收益率所绘制的频数多边形。

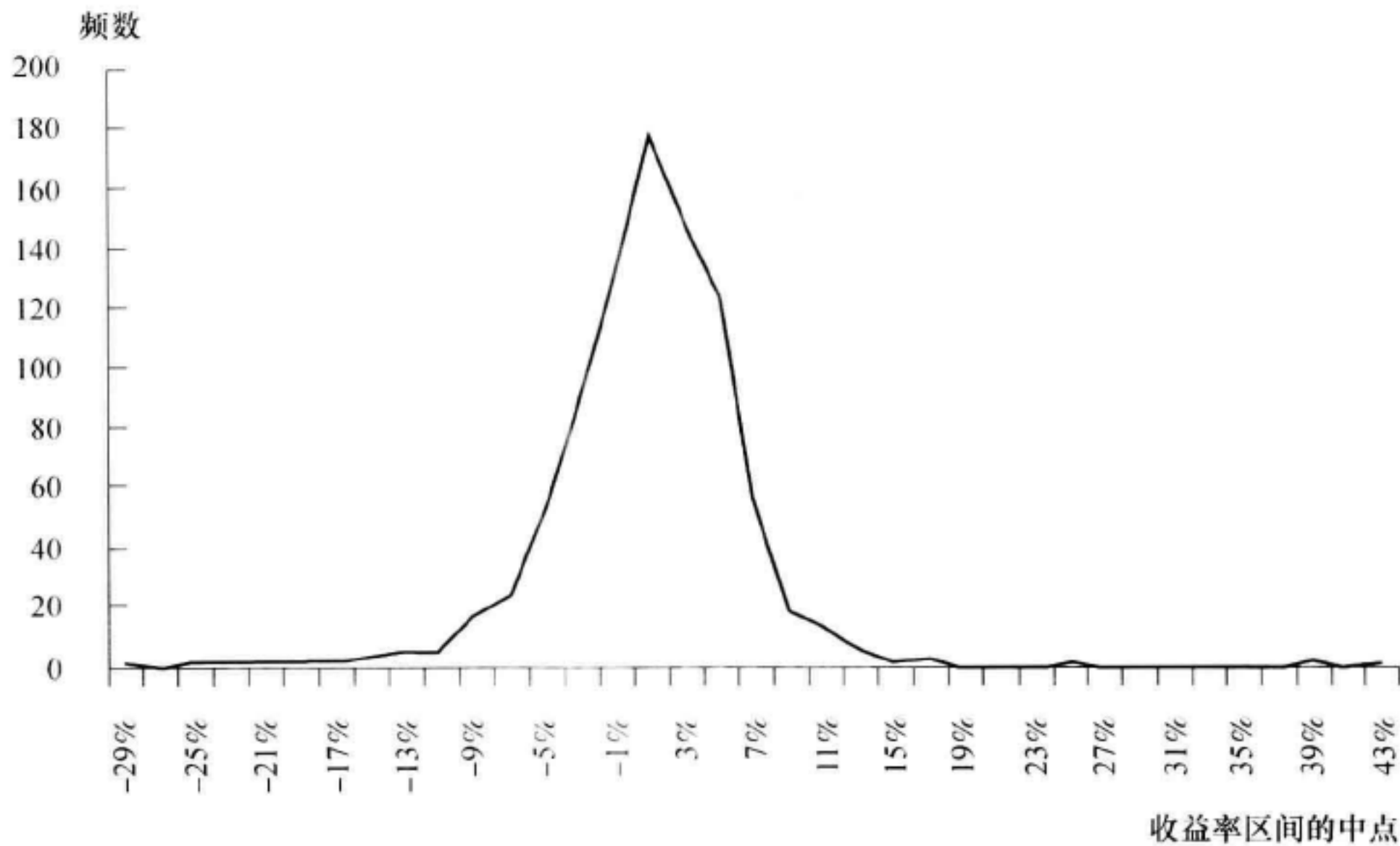


图 3-3 由 1926 年 1 月 ~ 2002 年 12 月标准普尔 500 月度总收益率所绘制的频数多边形
资料来源：伊博森协会。

在图 3-3 中，我们将直方图中的条状物用以直线相连的点来代替。例如 0% ~ 2% 的收益率区间的绝对频数为 182。在频数多边形中，我们画出了该收益率区间的中点为 1%，另外其频数为 182。我们按照这样类似的方法画出所有其他的点。[⊖]这种可视化的表示在分布描述过程中增加了一定程度的连续性特征。

另外一种形式的直线图是累积频数分布图。这类图可以将各个区间的累积绝对频数或者累积相对频数标注在图中的相应区间上限之上。累积频数分布图使得我们可以了解有多少观测值或者有多少百分比的观测值落在某一特定值的下方。为了制作累积频数分布，我们将表 3-4 中第 4 列或第 5 列中的数据绘制在每个收益率区间上限之上。图 3-4 显示的是由标准普尔 500 月收益率所绘制的累积绝对频数分布图。注意到当收益率为极小的负值或者极大的正值时，累积频数分布图倾向于变得平坦。图 3-4 正中央的陡峭部分反映出大部分观测值位于 -2% ~ 6% 区间之内的事实。

我们能进一步通过观察表 3-7 中所列出的两个收益率区间的情况来了解相对频数和累积相对频数之间的关系。第 1 个收益率区间 (0% ~ 2%) 的累积相对频数为 58.12%。第 2 个收益率区间 (2% ~ 4%) 的累积相对频数为 74.68%。在我们从一个区间转到另一个区间时，其累积相对频数产生了变化。例如，当我们从第 1 个收益率区间 (0% ~ 2%) 移到第 2 个区间 (2% ~ 4%) 时，累积相对频数的变化为 $74.68\% - 58.12\% = 15.56\%$ 。（表中的值都保留到小数点后两位。）

⊖ 即使区间的上界所代表的收益率并没有落入该区间，我们仍旧将其与区间下界进行平均以确定中点值。

其斜率是陡峭的事实表明这些频数是较大的。正如你可以从累积频数分布图中所看到的，曲线的斜率会随着我们从第一个收益率区间移向最后一个收益率区间而不断地发生变化。累积分布图中最前面的几个收益率区间的斜率相对较小，这告诉我们这些收益率区间没有包含很多的观测值。你可以回到表 3-4 中的频数分布中来加以确认，到第 10 个收益率区间（-12% ~ -10%）的累积绝对频数仅为 23 个观测值。大体上，累积绝对频数分布在任何特定区间的斜率与那个区间观测值的数目是成正比的。

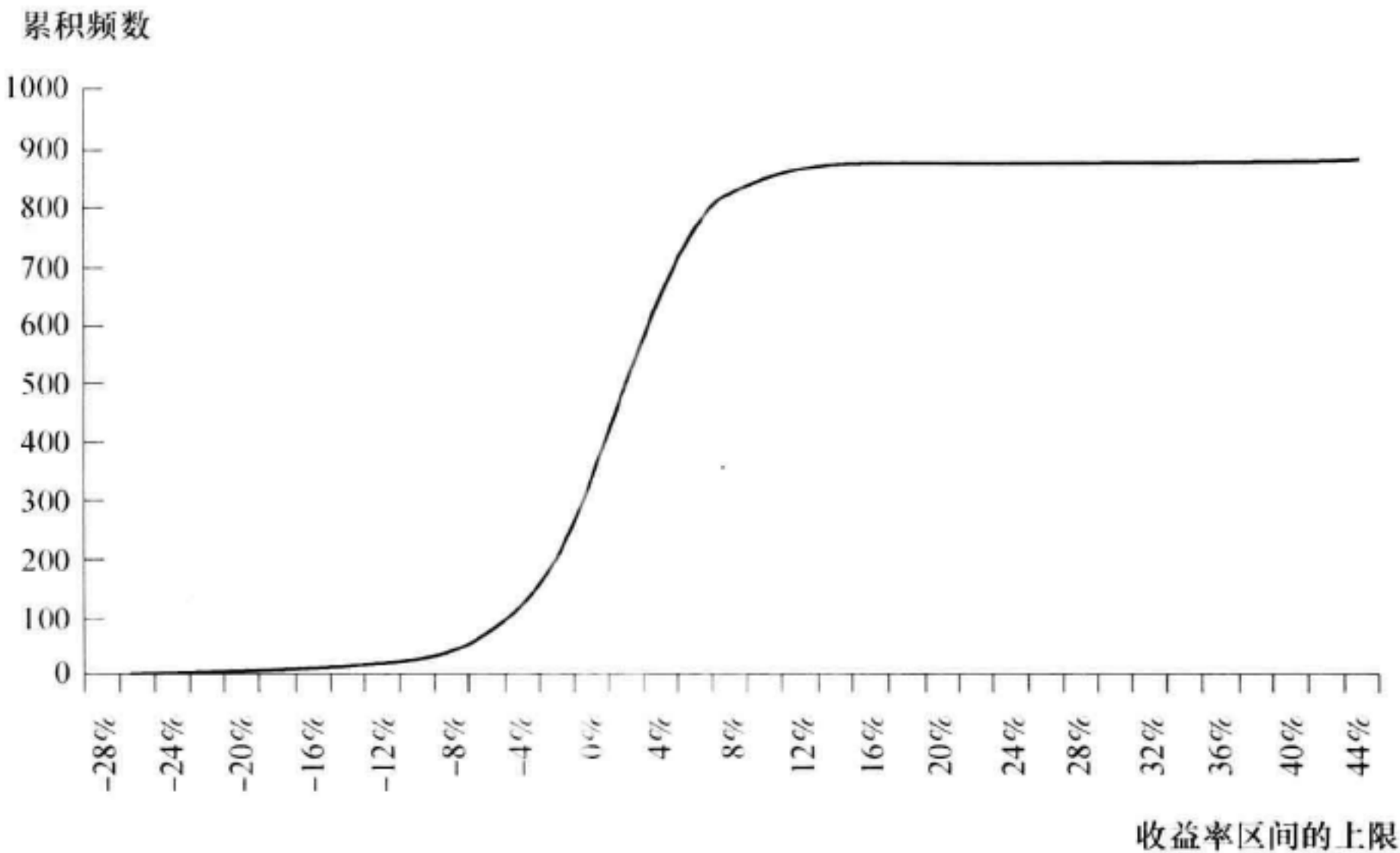


图 3-4 由 1926 年 1 月至 2002 年 12 月标准普尔 500 月度总收益率所绘制的累积绝对频数分布图
资料来源：伊博森协会。

表 3-7 选出的两个标准普尔 500 月收益率区间的频数

收益率区间 (%)	频数	相对频数 (%)	累积频数	累积相对频数 (%)
0.0 ~ 2.0	182	19.70	537	58.12
2.0 ~ 4.0	153	16.56	690	74.68

3.5 集中趋势的度量

目前，我们已经讨论了一些我们可以使用的、组织与表现数据以使其更加易于理解的方法。例如，一个资产类的收益率序列的频数分布揭示了投资者投资该资产类时可能面临的风险。正如上述图形中所描述的，标准普尔 500 年收益率的直方图明显地表明较大的正的年收益率和负的年收益率是非常普遍的。虽然频数分布和直方图提供了一个简便的汇总一系列观测值的方法，但是，这些方法只是描述数据的第一步。在这一节中，我们会讨论如何使用定量化度量指标来解释数据特征的方法。我们的重点将放在集中趋势测度和其他位置测度（或者位置参数）之上。集中趋势测度（measure of central tendency）详细说明了数据集中于何处。因为集中趋势测度容易计算和使用，所以相比其他的统计度量指标，其应用更加广泛。位置测度（measures of location）不仅包括了集中趋势测度，还包括了其他描述数据位置或分布的度量指标。

在下面的几个子小节中，我们会介绍一些常见的度量集中趋势的指标——算术平均值、中位数、众数、加权平均数和几何平均数。我们也会介绍其他一些有用的度量位置的指标，包括四分位数、五分位数、十分位数和百分位数。

3.5.1 算术平均数

分析师和投资组合经理经常想要获得一个能够描述投资决策代表性可能结果的数值。算术平均数是目前最常用的、用来度量数据中点或者中心的一个指标。

- **算术平均数的定义。**算术平均数 (arithmetic mean) 是所有观测值的总和除以观测值的数量后所得到的值。

我们可以计算出总体和样本的算术平均数，并分别称为总体均值和样本均值。

3.5.1.1 总体均值

总体均值是由一个总体所计算出的算术平均数。如果我们能完全地定义出一个总体，那么我们就可以计算出总体均值，即求出总体中所有观测值或数值的算术平均数。例如，调查 2002 财政年度美国主要零售俱乐部同类店铺销售增长情况的分析师可能会将所关注的总体限定在 3 家公司：BJ 零售俱乐部 (NYSE: BJ)，Costco 零售公司 (Nasdaq: COST) 和 Sam 俱乐部，沃尔玛公司的一部分 (NYSE: WMT)。[⊖]作为另一个例子，如果一位投资组合经理的投资集（必须从中选择的证券集合）为日经 - 道琼斯平均指数中的股票，那么相关的总体就是东京股票交易市场组成日经指数一板市场的 225 只股票。

- **总体均值计算公式 (population mean formula)。**总体均值 (population mean) μ 是总体的算术平均值。对于有限总体，总体均值为

$$\mu = \frac{\sum_{i=1}^N X_i}{N} \quad (3-2)$$

式中， N 表示整个总体中观测值的数量，而 X_i 表示第 i 个观测值。

总体均值是一个参数的例子。总体均值是唯一的，即一个给定的总体只会拥有一个均值。为了说明其计算方法，我们以在 2003 年 9 月初美国主要经营零售业这些公司的股票市盈率 (P/E) 的总体均值作为一个例子。在那一天，根据汤姆森金融 (Thomson Financial) First Call 数据库的数据，BJ、COST 和 WMT 的市盈率分别为 16.73、22.02 和 29.30。因此，在那一天市盈率的总体均值为 $\mu = (16.73 + 22.02 + 29.30)/3 = 68.05/3 = 22.68$ 。

3.5.1.2 样本均值

样本均值是由一个样本所计算出的算术平均数。许多时候我们并不能观测到集合中的每个元素；通常我们只能观测到总体的一个子集或者总体中的一个样本。总体均值的概念可以通过对于记号的一些修改，同样地应用到样本中的观测值之上。

- **样本均值计算公式 (sample mean formula)。**样本均值 (sample mean) 或者平均数， $\bar{X}$ (读做 “X-bar”) 是样本的算术平均值：

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n} \quad (3-3)$$

式中， n 表示样本中观测值的数量。

⊖ 一个零售俱乐部以一个商铺的形式运行，主要致力于将以仓库规模储存的大宗货物销售给支付它们会员费的客户。在 21 世纪早期，这 3 家零售俱乐部占据了美国零售业市场的绝大部分份额。

式 (3-3) 告诉我们将观测值 (X_i) 进行加总, 再除以观测值的数量 (就得到了样本均值)。例如, 如果样本的市盈率包含 35、30、22、18、15 和 12 这几个值, 那么市盈率的样本均值就是 $132/6 = 22$ 。样本均值同样被称为算术平均数。^①正如我们之前所讨论的, 样本均值是一个统计量 (即对于样本的描述性度量指标)。

均值可以对于不同个体或不同时间的数据进行计算。例如样本为金融时报交易所欧洲前 300 家公司指数 (该指数是由欧洲最大的 300 家公司股票组成的) 在 2003 年的净资产收益率。在这个例子中, 我们计算的是在 2003 年这 300 个个体净资产收益率的均值。当我们在一个特定的时间考察某些个体的特征时 (如金融时报交易所欧洲前 300 家公司指数的净资产收益率), 我们是在考察一组截面数据 (cross-sectional data)。这些观测值的均值被称为截面均值。而另一方面, 如果我们的样本包括的是金融时报交易所欧洲前 300 家公司指数在过去 5 年中的历史月收益率时, 那么我们所考察的就是时间序列数据 (time-series data)。这些观测值的均值被称为时间序列均值。我们将在时间序列分析的章节中讨论与时间序列行为相关的专门的统计方法。

下面, 我们要通过一个求解 2002 年 16 个欧洲国家的股票收益率样本均值的例子来进行说明。在这个例子中, 因为我们是将每个国家的收益率进行平均, 所以均值是截面的。

例 3-3 计算截面均值

摩根士丹利欧澳亚远东指数 (The MSCI EAFE (Europe, Australasia, and Far East)) 是一个经过浮动调整后的市值加权指数, 它是用来度量除美国和加拿大之外的成熟 (发达) 市场股票表现的一个指数。^②在 2002 年年末, 该指数包括 21 个发达市场国家的指数, 其中有 16 个欧洲市场、2 个澳洲市场 (澳大利亚和新西兰) 和 3 个远东市场 (中国香港、日本和新加坡)。

假设我们想了解 16 个欧洲市场在 2002 年 (大熊市的一年) 欧澳亚远东指数中以本国货币计量的表现情况。我们可以计算 2002 年这 16 个市场收益率的样本均值。表 3-8 中所给出的是以本国货币 (即用投资者所在国的货币计量所得到的收益率) 所度量的收益率序列。因为这个收益率不是以单个投资者所在国的货币所计量的, 因此, 它不是任何一个投资者所能实际获得的收益。事实上, 它是 16 个国家货币的加权平均收益率。

表 3-8 欧洲股票的总收益率

市场	以本国货币所计量的总收益率 (%)	市场	以本国货币所计量的总收益率 (%)
奥地利	-2.97	意大利	-23.64
比利时	-29.71	荷兰	-34.27
丹麦	-29.67	挪威	-29.73
芬兰	-41.65	葡萄牙	-28.29
法国	-33.99	西班牙	-29.47
德国	-44.05	瑞典	-43.07
希腊	-39.06	瑞士	-25.84
爱尔兰	-38.97	英国	-25.66

资料来源: www.mscidata.com.

① 相比“平均值”, 统计学家更喜欢用术语“均值”。一些作者将所有度量集中趋势的指标称为平均数 (包括中位数和众数)。而“均值”这个术语避免了任何可能造成的误解。
② 术语“自由浮动调整”指的是指数中所给予每个国家的权重反映的是那个国家实际能够进行投资的股票的价值。

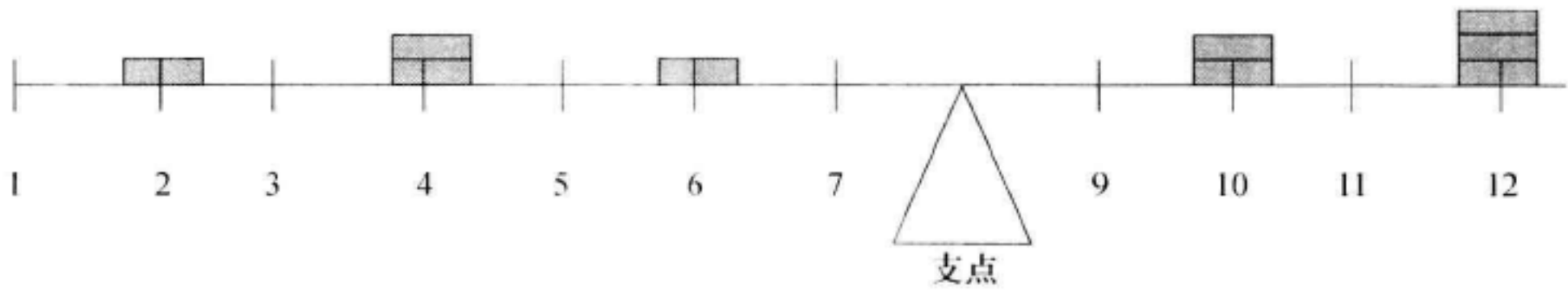
利用表 3-8 中的数据，我们可以计算出 2002 年这 16 只股票市场收益率的样本均值。

解：通过将表 3-8 所提供的收益率数据代入式 (3-3) 中，我们可以计算得到： $(-2.97 - 29.71 - 29.67 - 41.65 - 33.99 - 44.05 - 39.06 - 38.97 - 23.64 - 34.27 - 29.73 - 28.29 - 29.47 - 43.07 - 25.84 - 25.66)/16 = -500.04/16 = -31.25\%$ 。

在例 3-3 中，我们可以看到，有 7 个市场的收益率要低于均值，而另外 9 个市场的收益率高于均值。我们不应该期望任何实际的观测值与均值完全相等，因为样本均值提供的仅仅是对于所要分析数据的一个汇总。作为一位分析师，你将经常需要求解一些用以描述分布特征的数字。均值是一个常见的用以度量分布典型结果的统计量。你可以使用均值来比较两个不同市场间的表现情况。例如，你可能对于泛太平洋国家和欧洲国家的股票市场间的投资业绩比较感兴趣，这时你就可以使用这些市场的均值收益率来与投资结果进行比较。

3.5.1.3 算术平均数的性质

算术平均数可以比做一个物体的重心。图 3-5 通过将 9 个假设的观测值画在一条杆上来对于该类比给出一个图形上的表示。这 9 个观测值分别为 2, 4, 4, 6, 10, 10, 12, 12 和 12；算术平均值为 $72/9 = 8$ 。杆上所给出的观测值根据其频数以不同的高度来反映。（即 2 为 1 单位的高度，4 是 2 单位的高度，以此类推。）当该杆被放置到一个支点之上时，它只有在该支点正好被置于和其算术平均值相同大小的这一点时才会达到平衡。



当支点放在8处时，该杆就达到了完全平衡。

图 3-5 将算术平均数类比于（一个物体的）重心

作为分析师，我们经常将均值收益率作为一种资产可能出现的典型结果的一个度量指标。然而，正如前面的例子中所反映的，一些结果会高于均值，而另一些会低于均值。我们可以计算均值和每种结果之间的距离，并将其称为偏差值。从数学上来看，所有偏差值的总和总是等于 0 的。我们可以通过算术平均数的定义来揭示出这一点。

将式 (3-3) 两边同时乘以 n ： $n\bar{X} = \sum_{i=1}^n X_i$ 。于是，偏差值的总和就可以计算如下：

$$\sum_{i=1}^n (X_i - \bar{X}) = \sum_{i=1}^n X_i - \sum_{i=1}^n \bar{X} = \sum_{i=1}^n X_i - n\bar{X} = 0$$

算术平均值的偏差值是非常重要的信息，因为它可以反映风险。偏离均值的概念为更加复杂的方差、偏度和峰度概念提供了基础。我们将在本章的后面部分对这些更加复杂的概念进行讨论。

相对于其他两个集中度的度量指标（中位数和众数），算术平均数的优点在于它利用了所有观测值规模的信息。另外，均值很容易通过数学计算得到。

算术平均数的一个潜在缺点（也是它的一个性质）是它易于受到极端值的影响。因为所有观测值都被用来计算均值，所以算术平均数可能由于极端大或者极端小的观测值而产生严重的向上或向下的偏离。例如，假设我们计算下面 7 个数的算术平均数：1, 2, 3, 4, 5, 6 和 1 000。其

均值为 $1\,021/7 = 145.86$ 或者近似为 146。因为均值的大小 146 远远大于观测值中的大部分数（前 6 个数），我们可能会质疑其反映数据位置能力的好坏。在实务中，虽然金融数据集中的极端值或者异常值（outlier）可能仅仅反映总体的一种罕见情况，但是它也可能反映我们在记录观测值时所产生的错误，或者是从与总体不同的集合中所产生的观测值。特别在后两种情况下，所得到的算术平均数就可能会给我们带来误导。在这些情况下，最常见的处理方法是给出数据的中位数，而不是均值。^①下面我们就来讨论中位数的概念。

3.5.2 中位数

第二个重要的用来度量集中趋势的指标是中位数。

- **中位数的定义。**中位数（median）是指将数据集从小到大（升序）排列或者从大到小（降序）排列后，处于中心位置的数。在一个样本容量 n 为奇数的样本中，中位数处于 $(n+1)/2$ 的位置。而在一个样本容量为偶数的样本中，我们定义中位数为处于第 $n/2$ 和 $(n+2)/2$ 位置的两个数的均值（即最中间两个数的均值）。^②

之前我们给出的 3 家零售俱乐部的当前市盈率分别为：16.73，22.02 和 29.30。由于这是一个样本容量为奇数的样本（ $n=3$ ），所以中位数处于第 $(n+1)/2 = 4/2 = 2$ 的位置。故市盈率的中位数为 22.02。市盈率的值 22.02 是处于观测值最中间的数：一个在它上面，而另一个在它下面。无论我们是对于样本容量为偶数还是奇数的样本进行计算，位于中位数之上和之下的观测值的数量总是相等的。一个分布只有一个中位数。

中位数的潜在优点在于，与均值不同，极端值对其不会产生影响。然而，中位数并没有使用所有观测值的信息；它只关注于排序后观测值相对位置的情况。另外，中位数的计算也非常复杂；要计算中位数，我们需要将观测值从小到大排序，并确定样本大小是偶数还是奇数，并在此基础上，应用不同的计算方法。数学家将该缺点表述为中位数比均值在数学上难以处理。

为了展示如何求解中位数，我们使用例 3-3 的数据，在表 3-9 中以升序排列出了 2002 年欧洲股票的总收益率。因为这个例子有 16 个观测值，中位数应该是排序数列中位于第 $16/2 = 8$ 和第 $18/2 = 9$ 位的两个数的平均值。挪威的收益率 -29.73% 排在第 8 个位置，比利时的收益率 -29.71% 排在第 9 个位置。因此，由这两个收益率的均值所得到的中位数为： $(-29.73 - 29.71)/2 = -29.72\%$ 。值得注意的是，这里的中位数没有受到极端大或极端小的结果的影响。假设德国的总收益率改为一个更小的值或者是奥地利的总收益率变为一个更大的值，该数据的中位数并不会因此而发生改变。

① 其他处理极端值的方法都涉及对算术平均数计算的变形。截尾均值（Trimmed mean）在计算之前，先将最小与最大观测值中给定较小百分比的一部分舍去，然后再计算剩余值的算术平均数。例如，5% 的截尾均值将丢弃最小和最大各 2.5% 的数据，然后再计算剩余 95% 数据的算术平均数。截尾均值常被用于体育比赛之中，它通过去除裁判的最低分和最高分之后，给予参赛者一个平均得分。缩尾均值或温塞平均数（Winsorized mean）在计算均值前，先将数据中最小的一定百分比的值指定为一个相同的给定（低）数值，同时，将数据中最大的一定百分比的值指定为一个相同的给定（高）数值，然后再计算经过修改后的数据的平均数。例如 95% 的所谓均值将最小的 2.5% 的值都变为第 2.5 个百分位数的值，而将最大的 2.5% 的值都变为第 97.5 个百分位数的值。（我们将在后面给出百分位数的定义。）

② 记号 M_d 经常被用来表示中位数。和均值一样，我们需要区分总体中位数和样本中位数。了解到总体中位数将总体分成两半，而样本中位数将样本分成两半。为了简化，我们统一使用“中位数”这个术语来进行表述。

表 3-9 2002 年欧洲股票市场的总收益率（按升序排列）

编号	市场	以本币计算的总收益率（%）	编号	市场	以本币计算的总收益率（%）
1	德国	-44.05	9	比利时	-29.71
2	瑞典	-43.07	10	丹麦	-29.67
3	芬兰	-41.65	11	西班牙	-29.47
4	希腊	-39.06	12	葡萄牙	-28.29
5	爱尔兰	-38.97	13	瑞士	-25.84
6	荷兰	-34.27	14	英国	-25.66
7	法国	-33.99	15	意大利	-23.64
8	挪威	-29.73	16	奥地利	-2.97

资料来源：www.msdata.com.

例 3-4 在实际应用的背景下向我们展现了如何在存在极端值的样本下使用均值和中位数。

例 3-4 中位数和算术平均数：市盈率

假设一位客户要求你分析与评估一下表 3-10 中所给出的 7 个美国普通股投资组合的业绩情况。投资组合中的股票是用等权重加权的。你使用的评估度量指标是市盈率，即股票价格和每股收益的比率。虽然市盈率分母有许多不同算法，但是你所使用的市盈率指标定义为当前股价除以所有分析师对当前财政年度公司所估计出的每股收益的平均值。（列在“当前每股收益一致看法”一栏。）[⊖]表 3-10 中给出的是 2003 年 9 月 11 日的数值。为了进行比较，在那个时刻分析师对于标准普尔 500 当前市盈率的一致看法为 23.63。

表 3-10 客户投资组合的市盈率

股票	当前每股收益的一致看法	当前市盈率的一致看法
Exponent Inc. (Nasdaq: EXPO)	1.23	13.68
Express Scripts (Nasdaq: ESRX)	3.19	19.07
General Dynamics (NYSE: GD)	4.95	17.56
Limited Brands (NYSE: LTD)	1.06	15.60
Merant plc (Nasdaq: MRNT)	0.03	443.33
Microsoft Corporation (Nasdaq: MSFT)	1.11	25.61
O'Reilly Automotive, Inc. (Nasdaq: ORLY)	1.84	21.01

资料来源：First Call/Thomson Financial.

利用表 3-10 中的数据，给出下述问题的结果：

- (1) 计算市盈率的算术平均数。
- (2) 计算市盈率的中位数。
- (3) 通过市盈率的均值和中位数这两个度量集中趋势的指标评价上述投资组合的业绩。

解：

(1) 市盈率均值为 $(13.68 + 19.07 + 17.56 + 15.60 + 443.33 + 25.61 + 21.01) / 7 = 555.86 / 7 = 79.41$ 。

(2) 市盈率以升序排列为：

13.68 15.60 17.56 19.07 21.01 25.61 443.33

⊖ 关于价格乘数的更多内容，参见 Stowe, Robinson, Pinto 和 McLeavey (2002)。

样本含有奇数个观测值, $n = 7$, 所以中位数排在升序数列的第 $(n + 1)/2 = 8/2 = 4$ 个位置。因此, 市盈率中位数为 19.07。

(3) Merant 公司的市盈率约为 443, 其对于投资组合市盈率算术平均数的影响很大。市盈率均值 79 远远大于投资组合中 7 只股票中 6 只的市盈率。市盈率均值使我们错误地认为该投资组合具有很高的市盈率。排除 Merant 股票后的市盈率均值, 或者除去最大和最小股票的市盈率均值 (Merant 和 Exponent) 低于标准普尔 500 的市盈率 23.63。故市盈率的中位数 19.07 似乎是一个比较好的反映市盈率集中趋势的指标。

当公司的每股收益接近于 0 时——例如, 在一个商业周期的低点, 其市盈率将会变得异常的高。这一现象是经常出现的。在这种情况下的高市盈率反映了对于未来收益能力恢复的一种预期。对于极端的市盈率需要认真调查, 并小心处理。因为正如这个例子中所反映的, 分析师经常使用关于价格的比率指标的中位数来描述行业组类的估值特征。

3.5.3 众数

第 3 个重要的度量集中趋势的指标是众数。

- **众数的定义。**众数 (mode) 是指一个分布中发生频数最高的数值。^①

一个分布可能含有不止一个众数, 或者甚至没有众数。当一个分布只含有一个最频繁出现的数值时, 该分布被称为单峰的 (unimodal)。如果一个分布含有两个最频繁出现的数值时, 那么它就有两个众数, 并且该分布被称为双峰的 (bimodal)。如果一个分布含有三个最频繁出现的数值时, 那么它就是三峰的 (trimodal)。当数据集中所有的数值都不相同时, 该分布就没有众数, 因为没有数值发生的频率比任何其他数值更高。

股票收益率数据和其他连续型分布的数据可能没有众数。然而, 当这类数据被分组到区间中时, 我们经常可以找到一个区间 (可能超过一个区间) 拥有最高的发生频率: 众数区间 (modal interval) (或一组区间)。例如, 标准普尔 500 月收益率的频数分布拥有一个 0% ~ 2% 的众数区间, 如图 3-2 所示; 这个收益率区间拥有总体 924 个观测值中的 182 个观测值。众数区间总是直方图中最高的那个条状物所在的区间。

众数是集中趋势度量指标中唯一可以用于名义型数据的指标。当我们将共同基金分成不同的风格并对于其中的每种风格分配一个数字时, 这些分类数据的众数就是出现频率最高的共同基金的风格。

例 3-5 众数的计算

表 3-11 给出了 2002 年 9 月由穆迪的投资者服务部门所评级的 9 家美国百货公司发行的优先无担保债券的应用评级列表。按信用评级的降序排列 (或期望违约概率的升序排列), 穆迪评级依次为 Aaa, Aa1, Aa2, Aa3, A1, A2, A3, Baa1, Baa2, Baa3, Ba1, Ba2, Ba3, B1, B2, B3, Caa, Ca 和 C。^②

① 记号 M_o 通常用来表示众数。和均值和中位数一样, 我们可以对总体众数和样本众数进行区分。但是当我们了解了总体众数是总体中最大可能发生的数值, 而样本众数是样本中最大可能发生的数值之后, 为了简化, 我们一般统一采用术语“众数”来进行表述。

② 关于信用风险和信用评级的更多信息, 参见 Fabozzi (2004a)。

表 3-11 2002 年 9 月美国百货公司发行的优先无担保债券的评级

公司	信用评级	公司	信用评级
Dillards, Inc.	Ba3	Nordstom, Inc.	Baa1
Federated Department Stores, Inc.	Baa1	Penney, JC, Company, Inc.	Ba2
Kohl's Corporation	A3	Saks Incorporated	B1
May's Department Stores Company	A2	Sears, Roebuck and Co.	Baa1
Neiman Marcus Group, Inc.	Baa2		

资料来源：穆迪的投资者服务部门（Moody's Investors Service）。

利用表 3-11 中的数据，回答下列关于美国百货公司发行的优先无担保债券的相关问题：

- (1) 给出信用评级的众数。
- (2) 给出信用评级的中位数。

解：

(1) 这一组公司分别被给予了从 A2 ~ B1 的 7 种不同的信用评级。为了使我们的工作更加简单，我们首先将所有评级归整到一个频数分布中去。

表 3-12 美国百货公司发行的优先无担保债券信用评级的频数分布

信用评级	频数	信用评级	频数
A2	1	Ba2	1
A3	1	Ba3	1
Baa1	3	B1	1
Baa2	1		

除了 Baa2 的频数为 3，所有信用评级的频数都为 1。因此，根据那天穆迪的报告，美国百货公司的信用评级的众数应该为 Baa1。穆迪认为评级为 Baa1 的债券为中等风险的债务，它们既不受高度保护，也不是极度没有保障。

(2) 对于 $n = 9$ 的一组数据，其样本容量为奇数。该组数据的中位数位于第 $(n + 1)/2 = 10/2 = 5$ 的位置。我们可以从表 3-12 中看到，该值为 Baa1。因此，在 2002 年 9 月信用评级的中位数为 Baa1。

3. 5. 4 有关均值的其他概念

之前我们对于算术平均数进行了解释，它是描述数据集中趋势的一个基本概念。然而，其他一些有关均值的概念同样在投资中也非常重要。下面，我们就来对这些概念进行讨论。

3. 5. 4. 1 加权平均数

加权平均数的概念会多次出现在投资组合分析之中。对于算术平均数，所有观测值都是用相同的因子 $1/n$ 进行等值加权的。而在处理投资组合时，我们需要更加一般的加权平均的概念来使得对于不同的观测值拥有不同的权重。

为了解释加权平均数的概念，我们假设一名投资经理有一笔 100 000 000 美元需要投资，他可能将 70 000 000 美元投入股票，而将 30 000 000 美元投入债券。于是，投资组合有 0.70 的权重放在股票之上，而 0.30 的权重放在债券之上。那么，我们如何来计算该投资组合的投资收益率呢？投资收益率显然涉及对于股票和债券投资收益率的平均。然而我们所计算的均值必须反映投资组合中股票拥有 70% 的权重以及债券拥有的 30% 权重的这一事实。因此，反映这一权重的方法是

在股票投资收益率上乘以 0.70，而在债券投资收益率上乘以 0.30，然后将这两个结果进行加总。这一总和就是加权平均数。由股票收益率和债券收益率所算得的算术平均数，即给予两种资产相等权重的计算方法是不正确的。

考虑一个投资于加拿大股票和债券的投资组合，其中的股票是标准普尔/多伦多成分指数（S&P/TSX Composite Index）中的成分股，而债券是加拿大皇家银行资本市场的债券市场指数中的成分债券。这两个指数分别反映了加拿大股票和债券市场的总体情况。投资组合经理将投资组合中的 60% 分配于加拿大的股票，而将 40% 投资于加拿大的债券。表 3-13 给出了这两个指数从 1998 ~ 2002 年的总体收益率情况。

表 3-13 1998 ~ 2002 年加拿大股票和债券的总体收益率

年份	股票 (%)	债券 (%)
1998	-1.6	9.1
1999	31.7	-1.1
2000	7.4	10.3
2001	-12.6	8.0
2002	-12.4	8.7

资料来源：www.fidelity.ca 和 ww.money.msn.ca.

- 加权平均数公式（weighted mean formula）。对于一组观测值 $X_1, X_2, \dots, X_n$ 和对应的权重 $w_1, w_2, \dots, w_n$ ，加权平均数（weighted mean） $\bar{X}_w$ （读成“X-bar sub-w”）应该计算为：

$$\bar{X}_w = \sum_{i=1}^n w_i X_i \tag{3-4}$$

式中，权重之和等于 1；即 $\sum_{i=1}^n w_i = 1$ 。

在具体的投资组合中，正的权重表示买入并持有一项资产，而负的权重表示卖空一项持有的资产。[⊖]

我们所考虑的投资组合的收益率为加拿大股票和债券收益率的加权平均值（股票的权重为 0.60；债券的权重为 0.40）。除了成本，如果投资组合能够完全跟踪指数，那么我们利用式（3-4）可以计算得：

$$\begin{aligned} \text{1998 年的投资组合收益率} &= w_{\text{股票}} R_{\text{股票}} + w_{\text{债券}} R_{\text{债券}} \\ &= 0.60(-1.6\%) + 0.40(9.1\%) \\ &= 2.7\% \end{aligned}$$

我们应该清楚在该例子中所计算的正确均值应该是加权平均数，而不是算术平均数。假设我们计算了 1998 年的算术平均数，那么我们会计算得到其收益率为 $1/2(-1.6\%) + 1/2(9.1\%) = (-1.6\% + 9.1\%)/2 = 3.8\%$ 。已知投资组合经理 60% 投资于股票，40% 投资于债券。算术平均数会低估了投资于股票的权重，而高估了投资于债券的权重，结果所算得的投资组合的收益率数值要高出加权平均 1.1 个百分点（3.8% - 2.7%）。

现在假设该投资组合经理将在未来 5 年中保持 60% 投资于股票、40% 投资于债券的比例不变。这一投资方法被称为常比例投资策略。因为价值等于价格乘以数量，而价格波动将会引起权重的变化，所以常比例投资策略需要不断地将股票和债券的权重再次调整到所设定目标的水平。

⊖ 加权平均数的公式可以与算术平均数的公式相对比。如一组观测值 $X_1, X_2, \dots, X_n$ 和其对应的权重 $w_1, w_2, \dots, w_n$ （都等于 $1/n$ ），在这一假设下，加权平均数的公式将变为 $1/n \sum_{i=1}^n X_i$ 。这就是算术平均数的公式。因此，算术平均数是加权平均数在所有权重相等时的特例。

假设投资组合经理能够完成必要的再调整操作，我们可以利用式（3-4）计算出 1999 年、2000 年、2001 年和 2002 年的投资组合收益率，具体结果如下：

1999 年的投资组合收益率 = $0.60(31.7) + 0.40(-1.1) = 18.6\%$
2000 年的投资组合收益率 = $0.60(7.4) + 0.40(10.3) = 8.6\%$
2001 年的投资组合收益率 = $0.60(-12.6) + 0.40(8.0) = -4.4\%$
2002 年的投资组合收益率 = $0.60(-12.4) + 0.40(8.7) = -4.0\%$

我们现在可以利用计算算术平均数的式（3-3）来求解 1998 ~ 2002 年收益率时间序列的均值。投资组合总收益时间序列的均值为 $(2.7 + 18.6 + 8.6 - 4.4 - 4.0)/5 = 21.5/5 = 4.3\%$ 。

不同于计算投资组合年收益率时间序列的均值收益率，我们这里计算 5 年债券和股票收益率的算术平均数，然后分别给这两个收益率赋予 0.60 和 0.40 的权重。股票的平均收益率为 $(-1.6 + 31.7 + 7.4 - 12.6 - 12.4)/5 = 12.5/5 = 2.5\%$ 。债券平均收益率为 $(9.1 - 1.1 + 10.3 + 8.0 + 8.7)/5 = 35.0/5 = 7.0\%$ 。因此，投资组合的平均总收益率为 $0.60(2.5) + 0.40(7.0) = 4.3\%$ ，这与我们之前的计算是一致的。

例 3-6 作为加权平均数的投资组合收益率

表 3-14 给出了加拿大养老基金资产配置的平均估计收益率和 4 年期资产类收益率的信息。[⊖]

表 3-14 在 2003 年 3 月 31 日加拿大养老基金资产配置的平均估计收益率

资产类	资产配置 (权重)	资产类收益率 (%)	资产类	资产配置 (权重)	资产类收益率 (%)
股票	34.6	0.6	抵押贷款	1.3	9.0
美国股票	10.8	-9.3	房地产	4.5	10.2
国际股票	6.4	-10.5	现金及现金等价物	8.4	4.2
债券	34.0	6.0			

资料来源：标准生命投资有限公司（Standard Life Investments, Inc.）

利用表 3-14 中的信息，计算在 2003 年 3 月 31 日 4 年年末加拿大养老基金所获得的平均收益率。

解：将资产配置的百分比数据转化成小数点形式，我们可以计算出作为资产类收益率加权平均收益率的收益率均值。我们有：

投资组合平均收益率 = $0.346(0.6\%) + 0.108(-9.3\%) + 0.064(-10.5\%)$
 $+ 0.340(6.0\%) + 0.013(9.0\%) + 0.045(10.2\%)$
 $+ 0.084(4.2\%)$
 $= 0.208\% - 1.004\% - 0.672\% + 2.040\% + 0.117\%$
 $+ 0.459\% + 0.353\%$
 $= 1.5\%$

⊖ 在表 3-14 中，股票是以标准普尔/TSX 成份指数作为代表的，美国股票是以标准普尔 500 指数作为代表的，国际股票（非北美股票）是以 MSCI EAFE 指数为代表的，债券是以加拿大丰业银行资本市场（Scotia Capital Markets）统一债券指数为代表的，抵押贷款是以加拿大丰业银行资本市场（Scotia Capital Markets）抵押贷款指数为代表的，房地产是以标准生命投资房地产基金池为代表的，现金和现金等价物是以 91 天的国库券为代表的。

上面的例子向我们展现了一个一般的原则，即投资组合收益率是一个加权的总和。具体地，一个投资组合的收益率是投资组合中资产收益率的加权平均收益率；给予每个资产收益率的权重是其在投资组合中所投资资产的比例。

市场指数是以加权平均来计算的。对于市值加权指数如法国的 CAC-40 或者美国的标准普尔 500，每个指数中给予股票的权重等于其流通市值除以指数中所有股票的市值。

我们所介绍的都是利用历史数据所得到的加权平均数，但是加权平均数同样也可以利用预测数据来获得。当我们对于预测数据进行加权平均时，该加权平均数被称为期望值 (expected value)。假设我们要分别在经济扩张和经济收缩的情况下，对于年底标准普尔 500 的水平进行预测。如果我们分别赋予不同经济情况下的预测值以预测该情况发生的概率，将预测值与概率相乘，再将这两者加权预测相加，我们就计算出年末标准普尔 500 的期望值。如果我们求解的是标准普尔 500 未来可能收益率的加权平均数，那么我们就是在计算标准普尔 500 的期望值。其中，发生概率之和必须等于 1，另外，要满足式 (3-4) 加权平均数中权重表述的条件。

3.5.4.2 几何平均数

几何平均数是用来求解平均变化率或者变量增长率的最常用方法。在投资过程中，我们经常使用几何平均数来对于一个资产或者一个投资组合的收益率时间序列进行平均，或者对于金融变量如盈利或者销售额的增长率进行计算。例如，在货币时间价值一章中，我们计算过销售额的增长率（见例 1-17），而那个增长率是一个几何平均数。因为几何平均数十分重要，我们将在之后的小节中再次利用几何平均数，并提供一些关于其应用的实用观点。几何平均数是用如下的公式所定义的。

- **几何平均数公式 (geometric mean formula)**。对于一组观测值 $X_1, X_2, \dots, X_n$ 的几何平均数 (geometric mean) G 为

$$G = \sqrt[n]{X_1 X_2 X_3 \cdots X_n} \quad (3-5)$$

式中， $X_i \geq 0$ ，对于 $i = 0, 1, 2, \dots, n$ 。

只有当根号中乘积是非负的时候，式 (3-5) 才有解，或者说几何平均数存在。我们对于式 (3-5) 中的所有观测值 X_i 加以限制，令其大于等于 0。我们可以利用式 (3-5) 以及任何含有指数键的计算器（在大部分计算器上为 y^x ）来直接求解几何平均数。我们也能利用自然对数来求解几何平均数。式 (3-5) 也可以表示为：

$$\ln G = \frac{1}{n} \ln (X_1 X_2 X_3 \cdots X_n)$$

或者

$$\ln G = \frac{\sum_{i=1}^n \ln X_i}{n}$$

当我们计算出 $\ln G$ 之后， $G = e^{\ln G}$ （在大部分计算器中，这个步骤的按键为 e^x ）。

风险资产可能具有的最小负收益率为 -100%（如果它们的价格不低于 0 的话），因此我们必须在计算几何平均数时注意一下相关变量的定义。我们不能机械地将样本收益率进行连乘后，再对其求 n 次方根，因为在某个区间的收益率可能是负的。我们必须重新对于收益率进行定义，以使其都为正。我们通过在以小数形式表述的收益率基础上加上 1.0 来对其进行调整。 $(1 + R_t)$ 项代表相对于年初初始投资额的期末值增长倍数。只要我们使用 $(1 + R_t)$ ，那么观测值就不可能为

负，因为最大的负收益率为 -100% 。于是，该结果为 $1 + R_t$ 的几何平均数；再将所得结果减去 1，从而我们得到个体收益率 R_t 的几何平均数。例如，在表 3-13 中给出的 1998 ~ 2002 年间以标准普尔/TSX 成分指数表示的加拿大股票的收益率为 -0.016 ， 0.317 ， 0.074 ， -0.126 和 -0.124 （收益率以小数形式表示）。利用式（3-5），我们有 $\sqrt[5]{(0.9840)(1.317)(1.074)(0.874)(0.876)} = \sqrt[5]{1.065616} = 1.012792$ 。这个数是 1 加上收益率几何平均的值。从所得的结果中减去 1.0，我们得到 $1.012792 - 1.0 = 0.012792$ 或者约 1.3% 。在 1998 ~ 2002 年间加拿大股票收益率的几何平均数为 1.3% 。

一个汇总收益率几何平均数 R_G 的计算等式与之前的式（3-5）相比有稍稍的修改。其中 X_t 表达为“1 + 小数形式的收益率”。因为几何平均收益率使用的时间序列数据，我们同样使用 t 的下标。

$$1 + R_G = \sqrt[r]{(1 + R_1)(1 + R_2) \cdots (1 + R_r)}$$
$$1 + R_G = \left[\prod_{t=1}^r (1 + R_t) \right]^{\frac{1}{r}}$$

于是我们有下面的公式。

- 几何平均收益率公式 (geometric mean return formula)。给定一组持有期收益 R_t ， $t = 1, 2, \dots, T$ ，从 R_1 到 R_T 区间上的几何平均收益率为

$$R_G = \left[\prod_{t=1}^T (1 + R_t) \right]^{\frac{1}{T}} - 1 \tag{3-6}$$

我们可以使用式（3-6）来求解任何收益率数据序列的几何平均收益率。几何平均收益率也可以视为复合收益率。如果收益率是以月度频率用式（3-6）来平均的，那么我们可以将该月度几何平均收益率称为月度复合收益率。下面一个例子将向我们展现几何平均数的计算方法，同时反映几何平均数和算术平均数之间的差别。

例 3-7 几何平均和算术平均收益率（1）

作为一名共同基金分析师，你现在正在考察美国 2003 年年初时最近 5 年的两个大市值股票共同基金的总收益率。

表 3-15 1998 ~ 2002 年两家共同基金的总收益率

年份	所选的美国股票 (SLASX) (%)	普信股票收益基金 (PRFDX) (%)	年份	所选的美国股票 (SLASX) (%)	普信股票收益基金 (PRFDX) (%)
1998	16.2	9.2	2001	-11.1	1.6
1999	20.3	3.8	2002	-17.0	-13.0
2000	9.3	13.1			

资料来源：美国个人投资者协会 (American Association of Individual Investors, AAI)。

- 基于表 3-15 中的数据，回答下列问题：
- (1) 计算 SLASX 的几何平均收益率。
 - (2) 计算 SLASX 的算术平均收益率，并显示出其与基金几何平均收益率之间的差别。
 - (3) 计算 PRFDX 的几何平均收益率。
 - (4) 计算 PRFDX 的算术平均收益率，并显示出其与基金几何平均收益率之间的差别。

解：

(1) 将 SLASX 的收益率转化为小数形式，并将每个收益率加上 1.0，于是得：1.162，1.203，1.093，0.889 和 0.830。我们利用式 (3-6) 来求解 SLASX 的几何平均收益率：

$$\begin{aligned} R_G &= \sqrt[5]{(1.162)(1.203)(1.093)(0.889)(0.830)} - 1 \\ &= \sqrt[5]{1.127\,384} - 1 = 1.024\,270 - 1 = 0.024\,270 \\ &= 2.43\% \end{aligned}$$

(2) 对于 SLASX 来说， $\bar{R} = \frac{16.2 + 20.3 + 9.3 - 11.1 - 17.0}{5} = \frac{17.7}{5} = 3.54\%$ 。SLASX 的算术平均收益率超过几何平均收益率 $3.54 - 2.43 = 1.11\%$ 或者 111 个基点。

(3) 将 PRFDX 的收益率转化为小数形式，并将每个收益率加上 1.0，于是得：1.092，1.038，1.131，1.016 和 0.870。我们利用式 (3-6) 来求解 SLASX 的几何平均收益率：

$$\begin{aligned} R_G &= \sqrt[5]{(1.092)(1.038)(1.131)(1.016)(0.870)} - 1 \\ &= \sqrt[5]{1.133\,171} - 1 = 1.025\,319 - 1 = 0.025\,319 \\ &= 2.53\% \end{aligned}$$

(4) 对于 PRFDX 来说， $\bar{R} = \frac{9.2 + 3.8 + 13.1 + 1.6 - 13.0}{5} = \frac{14.7}{5} = 2.94\%$ 。PRFDX 的算术平均收益率超过几何平均收益率 $2.94 - 2.53 = 0.41\%$ 或者 41 个基点。表 3-16 汇总了所有的计算结果。

表 3-16 共同基金算术平均率和几何平均收益率：计算结果汇总

基金	算术平均数 (%)	几何平均数 (%)
SLASX	3.54	2.43
PRFDX	2.94	2.53

在例 3-7 中，对于这两个共同基金，几何平均收益率都低于算术平均收益率。事实上，几何平均数总是小于等于算术平均数。^①这两个平均数相等的唯一情况是在观测值不变的时候，即序列中所有的观测值都相同的时候。^②在例 3-7 中，由于基金收益率间存在不同，因此对于这两个基金的几何平均收益率都严格小于其算术平均收益率。通常，算术平均和几何平均间的差别会随着各期间观测值差别的增大而增大。^③这一关系同样可以从例 3-7 中看出。只要我们粗略地观察一下就会发现 SLASX 收益率间的差别要比 PRFDX 收益率间的差别大，因此，SLASX 算术平均收益率和几何平均收益率间的差别（111 个基点）要大于 PRFDX 算术平均收益率和几何平均收益率间的差别（41 个基点）。^④算术平均和几何平均也给出了这两个基金不同的排列顺序。虽然 SLASX 具有更高的算术平均收益率，但是 PRFDX 却具有更高的几何平均收益率。那么分析师是如何解释

① 这一结果可以利用詹森 (Jensen) 不等式来证明。如果函数从下向上看是凹的 (如 $\ln(X)$)，那么一个函数的平均值小于等于该函数在平均值处的值。

② 例如，假设 3 年中每年的收益率都为 10%。那么其算术平均数为 10%。为了求解几何平均数，我们先将收益率表示为 $(1 + R_t)$ 的形式，然后求解几何平均数如下： $[(1.10)(1.10)(1.10)]^{1/3} - 1.0 = 10\%$ 。这两个平均数是相等的。

③ 我们会马上介绍度量数据变动程度的指标——标准差。保持算术平均收益率不变，几何平均收益率会随着标准差的增加而减少。

④ 我们将在之后介绍度量数据变动程度的指标。但是我们要注意的，在 2000~2001 年间 SLASX 收益率的变动为 20.4%，而 PRFDX 收益率的变动为 11.5%。

这一结果的呢?

几何平均收益率表示的是一项投资的增长率或者是复合收益率。在 1998 年年初投资于 SLASX 的 1 美元将增长为 $(1.162)(1.203)(1.093)(0.889)(0.830) = \1.127 , 其值等于 1 加上几何平均收益率在 5 年中的复利值: $(1.0243)^5 = \$1.127$ 。这证实了几何平均数是复合收益率的结论。对于 PRFDX 来说, 1 美元将会增长到一个更大的值, $(1.092)(1.083)(1.131)(1.016)(0.870) = \1.133 , 等于 $(1.0253)^5$ 。如果我们把关注点放在多期投资的收益情况, 那么几何平均收益率是投资者首要关注的一个指标。当然, 算术平均收益率关注的是单期平均业绩的表现, 同样也是值得关注的。算术平均数和几何平均数都在投资管理中发挥着重要的作用, 对于一个收益率序列, 这两个指标通常都会在报告中给出。例 3-8 将之前的这些要点都在一个简单的例子中凸显出来。

例 3-8 几何平均收益率和算术平均收益率 (2)

假设一只股票的初始投资成本为 100 欧元。一年之后, 该股票交易价格为 200 欧元。在第 2 年年末, 该股票价格下降到原始的购买价格 100 欧元。在这两年间该股票没有进行红利支付。分别计算算术平均和几何平均年收益率。

解: 首先, 我们需要利用式 (3-1) 求解第 1 年和第 2 年的年收益率:

$$\text{第 1 年的收益率} = 200/100 - 1 = 100\%$$

$$\text{第 2 年的收益率} = 100/200 - 1 = -50\%$$

年收益率的算术平均数为 $(100\% - 50\%)/2 = 25\%$ 。

在我们求解几何平均数之前, 我们必须将百分比收益率转化为 $(1 + R_t)$ 。在进行调整之后, 由式 (3-6) 所得到的几何平均数为 $\sqrt{2.0 \times 0.50} - 1 = 0\%$ 。

0% 的几何平均收益率精确地反映了在第 2 年年末的投资价值等于第 1 年年初的投资价值。该投资的复利收益率为 0%。算术平均收益率反映的是一年期收益率的简单平均。

3.5.4.3 调和平均数

算术平均数、加权平均数和几何平均数是在投资中最常用的平均数的概念。第 4 个概念调和平均数 $\bar{X}_H$ 仅适合于有限种应用情形。[⊖]

- **调和平均数计算公式 (harmonic mean formula)**。对于一组观测值 $X_1, X_2, \dots, X_n$, 其调和平均数 (harmonic mean) 为:

$$\bar{X}_H = n / \sum_{i=1}^n (1/X_i) \quad (3-7)$$

式中, $X_i > 0$ 对于 $i = 1, 2, \dots, n$ 。

调和平均数是对观测值的倒数 (形式为 $1/X_i$ 的项) 进行求和后再求平均 (即将总和除以其观测值的数目 n), 最后将平均数求倒数后得到的。

调和平均数可以看做一种特殊形式的加权平均数, 其观测值的权重与其大小成反比例。调和平均数是关于均值相对专门的一个概念。当我们对比率进行平均 (每单位的数量) 且该比率被重复地应用于一个固定值来得到一个变化的单位数量时, 该调和平均数的计算是恰当的。这个概念

⊖ 术语“调和”来源于一种涉及倒数的序列, 其被称为调和序列。

最好通过一个例子来进行说明。一个广为人知的应用来源于一种投资策略，即成本平均，其涉及在每个区间中投资固定数目的一笔金额。在这个应用中，我们进行平均的比率是在购买当天的每股价格，而且我们将这些价格去除一个常数金额，并得出的一系列不同的股数值。

假设一个投资者在 $n=2$ 个月内每个月购买 1 000 欧元的股票，该股票的价格在这两个购买时点分别为 10 欧元和 15 欧元。那么购买该股票的平均价格是多少？

在这个例子中，第 1 个月我们购买了 $€1\,000/€10 = 100$ 股，第 2 个月我们购买了 $€1\,000/€15 = 66.67$ 或者总共 166.67 股。将总投资的欧元数额 $€2\,000$ 除以总购买的股数 166.67，所计算得到的平均购买价格为 $€2\,000/166.67 = €12$ 。该平均购买价格实际上就是关于各购买时点资产价格的调和平均数。利用式 (3-7)，调和平均价格为 $2/[(1/10) + (1/15)] = €12$ 。€12 这个值要小于购买价格的算术平均数 $(€10 + €15)/2 = €12.5$ 。然而我们可以利用加权平均公式发现 €12 才是正确的值，其购买价格的权重等于给定价格购买的股数相对于总购买股数的比率。在我们的例子中，该计算式为 $(100/166.67)€10.00 + (66.67/166.67)€15.00 = €12$ 。如果我们在每个时点投资不同的金额，那么我们就不能应用调和平均数的计算公式。但是，我们仍可以使用加权平均数的计算公式来进行类似于上面的计算。

一个关于调和平均数、几何平均数和算术平均数之间的数学关系是除非数据集的所有观测值都相同，否则调和平均数就将严格小于几何平均数，而几何平均数严格小于算术平均数。在我们给出的例子中，调和平均价格确实严格小于算术平均价格。

3.6 位置的度量：分位数

至今为止，我们已经讨论了集中趋势测度。现在我们来考察描述数据位置的一种方法，该方法被用来识别出这样的临界值，小于等于该临界值的观测值数量占到整个数据中某个具体比率。例如，一个投资组合 25%、50% 和 75% 的年收益率的值分别小于或者等于值 -0.05、0.16 和 0.25。这些数值提供了关于投资组合收益率分布的简要信息。统计学家使用词汇——分位数（或分位点）（quantile or fractile）来作为对于该值（小于等于该值的数据为一给定的比率）最一般的术语。下面，我们会讨论最常用的一些分位数——四分位数、五分位数、十分位数和百分位数，以及它们在投资中的一些应用。

3.6.1 四分位数、五分位数、十分位数、百分位数

我们知道中位数将一个分布分成相等的两个部分。我们可以定义其他将一个分布分成更小尺寸的分割线。四分位数（quartiles）将一个分布分成 4 等分，而五分位数（quintiles）将其分成 5 等分，十分位数（deciles）分成 10 等分，百分位数（percentiles）分成一百等分。给定一组观测值，第 y 个百分位数是指这样的值，小于等于该值的观测值占总观测值的比例为 $y\%$ 。百分位数经常被使用，其他的一些度量指标可以类似地由百分位点进行定义。例如，第 1 个四分位数（ Q_1 ）将一个分布分成两块，该分割点使得 25% 的观测值小于等于它。因此，第 1 个四分位点也就是第 25 个百分位点。同样的，第 2 个四分位点表示第 50 个百分位点，第 3 个四分位点表示第 75 个百分位点（因为 75% 的观测值小于等于它）。

当处理实际数据时，我们经常发现我们需要去近似估计百分位数的值。例如，如果我们对于第 75 个百分位数的值感兴趣，而且我们发现数据中没有使得样本中正好有 75% 的观测值小于等于它的观测值存在。我们就可以运用下面的方法来帮助我们确定或者估计该百分位数。这种方法

首先要求在该组观测值中确定百分位数的位置，然后确定（或者估计）出与那个位置相关的具体数值。

假使 P_y 是这样的一个数值，它使得分布的 $y\%$ 小于等于该值，或者说它就是第 y 个百分位点。（例如， P_{18} 就是使得分布中 18% 的观测值小于等于它的这样一个数值；或者说有 $100 - 18 = 82\%$ 的观测值比 P_{18} 大。）一个由 n 个按升序排列数据组成的数组的百分位点位置的计算公式为：

$$L_y = (n + 1) \frac{y}{100} \quad (3-8)$$

式中， y 是百分点（在该点我们将分布进行分割）； L_y 是按升序排序数组中位于百分数 P_y 处的值（ L ）。 L_y 该值可能是一个整数，也可能不是一个整数。通常，随着样本容量的增大，百分位数位置的计算将会变得越来越精确；而在小样本的情况下，它可能只是一个近似。

作为一个 L_y 不是整数的例子，假设我们要确定 2002 年由表 3-8 给出的 16 个欧洲股票市场收益率的第 3 个四分位数（ Q_3 或者 P_{75} ）。按照式（3-8），第 3 个四分位数的位置为 $L_{75} = (16 + 1) 75/100 = 12.75$ ，或者位于表 3-9 按收益率升序排列的第 12 项和第 13 项之间的值。表 3-9 中第 12 项为 2002 年葡萄牙的股票收益率 -28.29% 。第 13 项为 2002 年瑞士的股票收益率 -25.84% 。为了反映“12.75”中的“0.75”，我们得出这样的结论： P_{75} 位于 -28.29% 和 -25.84% 之间 75% 的距离上。

总结如下：

- 当位置 L_y 是一个整数时，该位置与一个实际的观测值相对应。例如，如果意大利没有被包含在样本中，那么 $n + 1$ 将等于 16，于是 $L_{75} = 12$ ，故第 3 个四分位数将为 $P_{75} = X_{12}$ ，其中， X_i 是有在按升序排列的数据中的第 i 个（ $i = L_{75}$ ）位置上的观测值的数值（即 $P_{75} = -28.29$ ）。
- 当 L_y 不是一个整数时， L_y 位于两个与其最接近的整数之间（一个大于它，一个小于它），那么我们可以通过在这两个位置之间进行线性插值（linear interpolation）来确定 P_y 。插值指的是通过两个已知的且在所求未知量附近的数值（一个大于它，一个小于它）来对于未知量进行估计的一种方法；而术语“线性”指的是一种利用直线来进行的估计。回到计算股票收益率的 P_{75} 值的例子中，我们之前求得 $L_y = 12.75$ ；其相邻的比它小的整数为 12，而相邻的比它大的整数为 13。利用线性插值方法， $P_{75} \approx X_{12} + (12.75 - 12)(X_{13} - X_{12})$ 。如上所述，在第 12 个位置上的数值是葡萄牙的股票收益率，所以 $X_{12} = -28.29\%$ ；而 $X_{13} = -25.84\%$ ，代表着第 13 个位置上的瑞士的股票收益率。因此，我们的估计为 $P_{75} \approx X_{12} + (12.75 - 12)(X_{13} - X_{12}) = -28.29 + 0.75[-25.84 - (-28.29)] = -28.29 + 0.75(2.45) = -28.29 + 1.84 = -26.45\%$ 。用语言来表达就是， P_{75} 被包含在一个下界为 -28.29 、上界为 -25.84 的区间中。因为 $12.75 - 12 = 0.75$ ，所以利用线性插值，我们将从 -28.29% 至 -25.84% 之间距离的 75% 处的点作为我们 P_{75} 的估计值。当 L_y 不是一个整数时，我们采用这种方式：取包含 P_y 的最接近的位于 L_y 之上和之下的两个整数来确定所使用观测值的位置，然后对于该位置上的两个观测值利用线性插值来对于 P_y 进行估计。

例 3-9 举例说明了欧洲主要股票指数中成分股红利收益率的不同分位数的计算方法。

例 3-9 计算百分位数、四分位数和五分位数

DJ EuroSTOXX 50 是一个由欧洲最大的 50 家上市公司按市值加权所计算得到的指数。表 3-17 列出了该指数中 50 个成分股在 2003 年中期的红利收益率情况，数据按升序排列。

表 3-17 DJ EuroSTOXX 50 种成分股的红利收益率

编号	公司	红利收益率 (%)	编号	公司	红利收益率 (%)
1	阿斯利康	0.00	26	瑞士联合银行	2.65
2	英国石油公司	0.00	27	乐购	2.95
3	德国电信	0.00	28	道达尔	3.11
4	汇丰控股	0.00	29	葛兰素史克	3.31
5	瑞士信贷集团	0.26	30	英国电信集团	3.34
6	欧莱雅	1.09	31	联合利华	3.53
7	瑞士再保险	1.27	32	巴斯夫	3.59
8	罗氏控股	1.33	33	西班牙国际银行	3.66
9	慕尼黑再保险集团	1.36	34	毕尔巴鄂比斯开银行	3.67
10	忠利集团	1.39	35	迪阿吉奥	3.68
11	沃达丰集团	1.41	36	哈利法克斯苏格兰银行	3.78
12	家乐福	1.51	37	德国意昂集团	3.87
13	诺基亚	1.75	38	壳牌运输贸易公司	3.88
14	诺华	1.81	39	巴克莱银行	4.06
15	安联	1.92	40	荷兰皇家石油公司	4.27
16	荷兰皇家飞利浦电子	2.01	41	富通集团	4.28
17	西门子	2.16	42	拜耳	4.45
18	德意志银行	2.27	43	戴姆勒克莱斯勒	4.68
19	意大利电信	2.27	44	苏伊士集团	5.13
20	法国安盛	2.39	45	英杰华集团	5.15
21	西班牙电信集团	2.49	46	埃尼集团	5.66
22	雀巢	2.55	47	荷兰国际集团	6.16
23	苏格兰皇家银行集团	2.60	48	保诚集团	6.43
24	荷兰银行	2.65	49	骏懋（劳埃德）银行	7.68
25	法国巴黎银行	2.65	50	荷兰全球人寿保险公司	8.14

资料来源：<http://france.finance.yahoo.com>，2003 年 7 月 8 日。

利用表 3-17 中的数据回答下列问题：

- (1) 计算第 10 个和第 90 个百分位数。
- (2) 计算第 1 个、第 2 个和第 3 个四分位数。
- (3) 给出中位数的值。
- (4) 有多少个五分位数，这些五分位数分别与哪些百分位数相对应？
- (5) 计算第 1 个五分位数。

解：

(1) 在这个例子中， $n = 50$ 。利用式 (3-8)， $L_y = (n + 1)y/100$ 是第 y 个百分位数所处的位置，因此，对于第 10 个百分位数来说，我们有

$$L_{10} = (50 + 1)(10/100) = 5.1$$

L_{10} 在第 5 个和第 6 个观测值 $X_5 = 0.26$ 和 $X_6 = 1.09$ 之间。红利收益率的第 10 个百分位数 (第 1 个十分位数) 的估计值为

$$\begin{aligned} P_{10} &\approx X_5 + (5.1 - 5)(X_6 - X_5) = 0.26 + 0.1(1.09 - 0.26) \\ &= 0.26 + 0.1(0.83) = 0.34\% \end{aligned}$$

对于第 90 个百分位数,

$$L_{90} = (50 + 1)(90/100) = 45.9$$

L_{90} 位于第 45 个观测值 $X_{45} = 5.15$ 和第 46 个观测值 $X_{46} = 5.66$ 之间。故第 90 个百分位数 (第 9 个十分位数) 的估计值为

$$\begin{aligned} P_{90} &\approx X_{45} + (45.9 - 45)(X_{46} - X_{45}) = 5.15 + 0.9(5.66 - 5.15) \\ &= 5.15 + 0.9(0.51) = 5.61\% \end{aligned}$$

(2) 第 1 个、第 2 个和第 3 个四分位数分别对应于 P_{25} 、 P_{50} 和 P_{75} 。

$$\begin{aligned} L_{25} &= (51)(25/100) = 12.75 \quad L_{25} \text{ 位于第 12 和第 13 个数据} \\ &\quad X_{12} = 1.51 \text{ 和 } X_{13} = 1.75 \text{ 之间。} \end{aligned}$$

$$\begin{aligned} P_{25} &= Q_1 \approx X_{12} + (12.75 - 12)(X_{13} - X_{12}) \\ &= 1.51 + 0.75(1.75 - 1.51) \\ &= 1.51 + 0.75(0.24) = 1.69\% \end{aligned}$$

$$\begin{aligned} L_{50} &= (51)(50/100) = 25.5 \quad L_{50} \text{ 位于第 25 和第 26 个数据之间,但是} \\ &\quad \text{由于 } X_{25} = X_{26} = 2.65, \text{ 故我们不需要插值。} \end{aligned}$$

$$P_{50} = Q_2 = 2.65\%$$

$$\begin{aligned} L_{75} &= (51)(75/100) = 38.25 \quad L_{75} \text{ 位于第 38 和第 39 个数据} \\ &\quad X_{38} = 3.88 \text{ 和 } X_{39} = 4.06 \text{ 之间。} \end{aligned}$$

$$\begin{aligned} P_{75} &= Q_3 \approx X_{38} + (38.25 - 38)(X_{39} - X_{38}) \\ &= 3.88 + 0.25(4.06 - 3.88) \\ &= 3.88 + 0.25(0.18) = 3.93\% \end{aligned}$$

(3) 中位数是第 50 个百分位数, 2.65%。这与我们通过求第 $n/2 = 50/2 = 25$ 项和第 $(n+2)/2 = 52/2 = 26$ 项的平均数的结果是相同的。这一方法与之前样本容量为偶数的样本中位数的求解过程所得的结果是一致的。

(4) 有 4 个五分位点, 它们分别对应 P_{20} 、 P_{40} 、 P_{60} 和 P_{80} 。

(5) 第 1 个五分位点是 P_{20} 。

$$\begin{aligned} L_{20} &= (50 + 1)(20/100) = 10.2 \quad L_{20} \text{ 位于第 10 个和第 11 个观测值} \\ &\quad X_{10} = 1.39 \text{ 和 } X_{11} = 1.41 \text{ 之间} \end{aligned}$$

故第 1 个五分位数的估计值为

$$\begin{aligned} P_{20} &\approx X_{10} + (10.2 - 10)(X_{11} - X_{10}) \\ &= 1.39 + 0.2(1.41 - 1.39) \\ &= 1.39 + 0.2(0.02) = 1.394\% \text{ 或 } 1.39\% \end{aligned}$$

3.6.2 分位数在投资中的应用

在这一节中，我们将对分位数在投资中的应用进行讨论。分位数在投资组合业绩评价和投资策略研发中都有应用。

投资分析师每天使用分位数来对于业绩进行排序——例如，对于投资组合的业绩。投资经理的业绩也经常以与同等类型投资经理的业绩相比所处的分位数的形式来表示。另外，晨星共同基金星级评级也将同类风格共同基金百分位数的业绩与某个星级的数字相联系。

另外一个使用分位数的领域是在投资研究方面。分析师根据某一分位数来定义一个以该分位数命名的组类。例如，分析师们经常将一组收益率低于作为截断点的第10个百分位数的公司作为收益率最低的10%的公司。基于某些特征将数据分成几类，使得分析师们能评估特定特征对于某些感兴趣数据的影响情况。例如，金融实证研究常根据股票市值对上市公司进行排序并将其分成10类。第1类包含的是那些拥有最小市值公司投资组合，而第10类包含的是那些拥有最大市值公司的投资组合。将上市公司分成10类使得分析师们能够比较小市值公司和大市值公司业绩表现之间的差别。

我们可以用鲍曼（Bauman）、柯诺瓦（Conover）和米勒（Miller）（1998）作为一个在投资研究中使用分位数，特别是四分位数的例子来加以说明。他们的研究比较了国际成长股和价值股之间的业绩表现情况。通常，价值股被定义为那些市场价格相对其每股盈利、每股账面价值或者每股红利较低的股票。而另一方面，成长股的市场价格却相对于那些指标较高。鲍曼他们的分类标准是按照一些估值度量指标：价格－盈利比（市盈率 P/E ），价格－现金流比（ P/CF ），价格－账面价值比（市净率， P/B ）和红利收益率（ D/P ）。他们将总样本中在1986～1996年间每年6月20日拥有最低市盈率的四分之一股票（价值股类）定义为四分位数组1，而将每年拥有最高市盈率的四分之一股票（成长股类）定义为四分位数组2。拥有第二大的市盈率的股票形成四分位数组3，而拥有第二小的市盈率的股票形成四分位数组4。作者们对于每个基本面因子重复这个分类过程。对于每个四分位数组类，他们将其以等权重的方式形成一个投资组合。由此，他们能够比较不同价值/成长四分位数组类间的业绩表现情况。下面的表3-18是由他们研究中的表1复制而来的。

表 3-18 基于所选特征分类的价值股和成长股的平均年收益率（1986～1996 年）

选择标准	总观测数	Q_1 (价值型)	Q_2	Q_3	Q_4 (成长型)	Q_1 与 Q_4 间 收益率的差
按 P/E 分类	28 463					
P/E 中位数		8.7	15.2	24.2	72.5	
收益率		15.0%	13.6%	13.5%	10.6%	+4.4%
标准差		46.5	38.3	42.5	50.4	
按 P/CF 分类	30 240					
P/CF 中位数		4.4	8.2	13.3	34.2	
收益率		15.5%	13.7%	12.9%	11.2%	+4.3%
标准差		48.7	41.2	41.9	51.4	
按 P/B 分类	32 265					
P/B 中位数		0.8	1.4	2.2	4.3	
收益率		18.1%	14.4%	12.6%	12.4%	+5.7%
标准差		69.6	45.9	45.1	57.0	
按 D/P 分类	25 394					
D/P 中位数		5.6%	3.2%	1.9%	0.6%	
收益率		14.1%	14.1%	12.5%	9.3%	+4.8%
标准差		40.5	38.7	38.9	42.0	

资料来源：鲍曼（Bauman）等。

表 3-18 给出了每个四分位数组类以及每个估值因子的中位数、均值收益率和标准差。从四分位数组 1 到四分位数组 4，按 P/E、P/CF 和 P/B 分类时，都呈现逐渐增加的趋势；而按 D/P 分类时，呈现出逐渐下降的趋势。而无论选择标准是什么，国际价值股在样本区间内总是表现优于成长股。

鲍曼、柯诺瓦和米勒也根据股票市值的大小将公司分入 4 个四分位数组类之一。在此基础上，他们考察了在不同四分位数组类中股票的收益率情况。表 3-19 是从他们论文中的表 7 复制而来。如表中所示，小公司的投资组合具有市值中位数为 46 600 000 美元，而大公司的投资组合具有的市值中位数为 2 472 300 000 美元。大公司的市值是小公司的 50 多倍，但是它们的股票平均收益率却小于小公司平均收益率的一半（小公司 22.0%，大公司 10.8%）。总的来看，鲍曼等人发现两种效应：第一，国际价值股（如作者们所定义的）的表现要优于国际成长股；第二，国际小市值股票的表现要优于国际大市值的股票。

表 3-19 按市值分组的国际股票的平均年收益率（1986 ~ 1996 年）

分类标准	总观测值	Q_1 (小规模)	Q_2	Q_3	Q_4 (大规模)	Q_1 与 Q_4 间 收益率的差
按规模分类	32 555					
规模中位数（以百万美元计）		46.6	209.9	583.7	2 472.3	
收益率		22.0%	13.6%	11.1%	10.8%	+ 11.2%
标准差		87.8	45.2	39.5	34.0	

资料来源：鲍曼等。

作者下一步是考察在控制公司规模的情况下，价值股和成长股表现情况如何。这一步涉及构建 16 个不同的价值型/成长型和公司规模的投资组合（ $4 \times 4 = 16$ ），并观察这两个基本面因素相互间的交互作用。他们发现除了市值非常小的情况之外，国际价值股票的表现始终优于国际成长股。对于投资组合经理来说，这些发现告诉我们在哪个具体研究的时间区间内将国际市场上的价值股推荐给投资者比推荐成长股更能获得有利的收益率。

3.7 离散度的度量

正如著名的研究者费雪·布莱克（Fisher Black）所写的，“投资中的重要问题在于估计期望收益率”。^①很少有人会忽视投资中期望收益率或者平均收益率的重要性：平均收益率告诉我们收益率和投资的结果会集中出现在哪里。然而为了完全地理解投资，我们也需要去了解收益率在均值附近的分散情况。离散度（dispersion）是反映集中趋势附近变化程度的一个度量。如果平均收益率表示的是回报，那么离散度表达的就是风险。

在这一节中，我们将考察最常用的离散度度量指标：极差、平均绝对偏差、方差和标准差。所有这些都是度量绝对离散度（absolute dispersion）的指标。绝对离散度是指不需要与任何参考点或基准进行比较而得到的反映数据变动程度的量。

这些度量指标将在整个投资实践中得到使用。收益率的方差和标准差经常被用来作为风险的度量指标。这一方法最早起源于诺贝尔奖得主哈里·马克维茨（Harry Markowitz）的开创性工作。而另一位诺贝尔经济学奖得主威廉·夏普（William Sharpe）则提出了夏普比率，这是一个经过风险调整后的业绩度量指标。该指标利用了收益率的标准差作为风险的度量。其他度量数据离散度

① 布莱克（Black）（1993）。

的指标——平均绝对偏差和极差同样也在数据分析中非常有用。

3.7.1 极差

我们先前在讨论构建频数分布过程中就已经遇到过极差。作为所有度量离散度指标中最简单的指标，极差可以对区间型或者比率型数据进行计算。

- **极差的定义。**极差 (range) 是指一个数据集中最大值与最小值之差。

极差 = 最大值 - 最小值 (3-9)

作为极差的一个例子，从 1926 年 1 月至 2002 年 12 月标准普尔 500 最大的月收益率为 42.56%（出现在 1933 年 4 月），最小的收益率为 -29.73%（出现在 1931 年 9 月），因此收益率的极差为 72.29% [42.56% - (-29.73%)]。另外一种代替极差定义的方法是报告最大值和最小值。这种替代的方法能比式 (3-9) 的定义提供更多的信息。

极差的一个优点在于计算简单。而其缺点在于极差仅仅使用了分布中的两个信息，它无法告诉我们数据是如何分布的（即分布的形状是怎样的）。因为极差是最大值和最小值之差，它只能反映出极端大和极端小的情况，而不能表现出分布总体的情况。[⊖]

3.7.2 平均绝对偏差

离散度的度量指标可以利用分布中所有的观测值来进行计算，而不是仅仅使用最大和最小值。但问题是我们应该如何来度量离散度呢？我们最初关于算术平均数的讨论中介绍了离开均值的距离或偏差的概念 ($X_i - \bar{X}$)，并将其作为统计学中使用的一个基本信息。我们可以通过考察偏离均值偏差的算术平均来计算度量数据离散度的指标，但是在这过程中我们遇到一个问题：偏离均值的偏差之和总是等于 0。如果我们计算偏差的平均数的话，这个结果将同样为 0。因此，我们需要寻找一个将负的偏差值从正的偏差值中除去的方法。

一种方法为考察偏离均值的绝对偏差，以计算出平均绝对偏差值。

- **平均绝对偏差计算公式 (mean absolute deviation formula)。**对于一个样本的平均绝对偏差 (mean absolute deviation, MAD) 为：

$$MAD = \frac{\sum_{i=1}^n |X_i - \bar{X}|}{n}$$
 (3-10)

式中， $\bar{X}$ 表示样本均值； n 表示样本中观测值的数目。

在计算 MAD 中，我们忽略了偏离均值偏差的符号。例如，如果 $X_i = -11.0$ 且 $\bar{X} = 4.5$ ，差别的绝对值为 $|-11.0 - 4.5| = |-15.5| = 15.5$ 。平均绝对偏差由于利用样本中所有的观测值，所以其比极差更适合作为度量离散度的指标。MAD 的一个技术性缺陷在于它相对于我们下面介绍的度量指标——方差，在数学上难以处理。[⊖] 例 3-10 举例介绍了如何利用极差和平均绝对偏差

⊖ 我们可能遇到的另一个度量离散度的距离指标为四分位距，它关注于数据的中间部分，而不是极端值。四分位距 (interquartile range, IQR) 是数据集中第 3 个和第 1 个四分位数之差：IQR = $Q_3 - Q_1$ 。IQR 代表包括中间 50% 数据的区间长度。其他情况不变，四分位距越大则表示数据离散度越大。

⊖ 在一些分析性的工作中（如最优化），对于微分的可积运算是重要的。方差可以作为一个可微函数，而绝对值却不

来对于风险进行度量。

例 3-10 极差和平均绝对偏差

例 3-7 中我们已经计算出了两家共同基金的平均收益率。现在分析师们将对于它们的风险进行评估。

表 3-15 (续) 两家共同基金的总收益率 (1998 ~ 2002 年)

年份	所选的美国股票 (SLASX) (%)	T. Rowe 价格 股票收益 (%)	年份	所选的美国股票 (SLASX) (%)	T. Rowe 价格 股票收益 (%)
1998	16.2	9.2	2001	-11.1	1.6
1999	20.3	3.8	2002	-17.0	-13.0
2000	9.3	13.1			

资料来源: AAIL.

基于表 3-15 中所给出的数据, 回答下列问题:

(1) 计算 (A) SLASX 和 (B) PRFDX 年收益率的极差, 并基于所得极差说明哪一个共同基金的风险更大。

(2) 计算 (A) SLASX 和 (B) PRFDX 年收益率的平均绝对偏差, 并基于所得的 MAD 说明哪一个共同基金的风险更大。

解:

(1) A. 对于 SLASX, 最大收益率为 20.3%, 最小收益率为 -17.0%。因此, 极差为 $20.3 - (-17.0) = 37.3\%$ 。

B. 对于 PFRDX, 极差为 $13.1 - (-13.0) = 26.1\%$ 。由于与 PRFDX 相比, SLASX 有更大的收益率极差, 因此, SLASX 在 1998 ~ 2002 年间是风险更大的基金。

(2) A. 在例 3-7 中所计算出的 SLASX 算术平均收益率为 3.54%。故 SLASX 收益率的平均绝对偏差值为

$$\begin{aligned} \text{MAD} &= \frac{|16.2 - 3.54| + |20.3 - 3.54| + |9.3 - 3.54| + |-11.1 - 3.54| + |-17.0 - 3.54|}{5} \\ &= \frac{12.66 + 16.76 + 5.76 + 14.64 + 20.54}{5} \\ &= \frac{70.36}{5} = 14.1\% \end{aligned}$$

B. 在例 3-7 中所计算出的 PRFDX 算术平均收益率为 2.94%。故 PRFDX 收益率的平均绝对偏差值为

$$\begin{aligned} \text{MAD} &= \frac{|9.2 - 2.94| + |3.8 - 2.94| + |13.1 - 2.94| + |1.6 - 2.94| + |-13.0 - 2.94|}{5} \\ &= \frac{6.26 + 0.86 + 10.16 + 1.34 + 15.94}{5} \\ &= \frac{34.56}{5} = 6.9\% \end{aligned}$$

SLASX 的 MAD 为 14.1%, 故其要比具有 6.9% MAD 的 PRFDX 的风险更大。

3.7.3 总体方差和总体标准差

平均绝对偏差通过对于偏差值取绝对值再求和的方法,解决了偏离均值的偏差值之和为0的问题。而第2种解决该问题的方法是将偏差值进行平方。基于偏差值平方的方差和标准差就是两个应用最为广泛的度量离散度的指标。方差定义为一个随机变量与其均值偏差平方的平均值。标准差为方差正的平方根。下面我们对于方差和标准差的计算和应用进行讨论。

3.7.3.1 总体方差

如果我们知道总体中的每一个元素,那么我们就可以计算出总体方差的值。总体方差是偏离均值偏差平方的算术平均数,我们用符号 σ^2 来表示。

- **总体方差计算公式 (population variance formula)**。总体方差 (population variance) 为

$$\sigma^2 = \frac{\sum_{i=1}^N (X_i - \mu)^2}{N} \quad (3-11)$$

式中, μ 表示总体均值,而 N 表示总体的规模。

如果给定总体均值,我们可以利用式(3-11)来计算偏离均值偏差的平方和(这里,考虑了总体中所有 N 项),然后通过将平方和除以 N 来求得平均平方偏差。无论偏离均值的数值为正还是为负,偏差的平方总是一个正数。因此,方差通过对于偏差进行平方的操作,解决了偏离均值偏差中存在的负数的问题。之前给出的BJ、COST和WMT的市盈率分别为16.73、22.02和29.30。我们所算出的市盈率均值为22.68。因此,市盈率的总体方差为 $(1/3)[(16.73 - 22.68)^2 + (22.02 - 22.68)^2 + (29.30 - 22.68)^2] = (1/3)((-5.95)^2 + (-0.66)^2 + (6.62)^2) = (1/3)(35.4025 + 0.4356 + 43.8244) = (1/3)(79.6625) = 26.5542$ 。

3.7.3.2 总体标准差

因为方差是以单位的平方进行计量的,我们需要寻找一个将其回复到原来单位的方法。我们可以利用标准差,即方差的平方根来解决这个问题。标准差比起方差更加容易解释,因为标准差所用的度量单位与观测值是相同的。

- **总体标准差计算公式 (population standard deviation formula)**。总体标准差 (population standard deviation) 定义为总体方差的正的平方根,具体为:

$$\sigma = \sqrt{\frac{\sum_{i=1}^N (X_i - \mu)^2}{N}} \quad (3-12)$$

式中, μ 表示总体均值, N 表示总体的规模。

利用BJ、COST和WMT市盈率的例子,由式(3-12),我们可以得到方差为26.5542。接着我们对其求平方根: $\sqrt{26.5542} = 5.1531$ 或者约为5.2。

总体方差和总体标准差都是分布参数的例子。在后面的章节中,我们将会介绍作为风险度量指标的方差和标准差的概念。

在投资活动中,我们经常不知道所研究总体的均值,因为我们一般无法从总体中识别出每个元素或者对每一个元素进行计算。于是,我们需要通过从总体中抽取出的样本均值来对于总体均值进行估计。我们可以利用与式(3-11)和式(3-12)不同的公式来计算出样本方差和样本标准

差。我们将在之后的几节中对这些计算方法进行讨论。然而，在投资活动中，有时我们有一个定义良好的可以被视为总体的组类。对于良好定义的总体，我们就可以使用式（3-11）和式（3-12）来进行相应计算，具体如例 3-11 所示。

例 3-11 计算总体标准差

表 3-20 给出了组成 2002 年《福布斯》杂志荣誉榜的 10 家美国股票型基金的年度投资组合换手率。^①投资组合换手率是一个度量交易活动的指标，其值为一年中销售值和购买值中的最小者除以该年的平均净资产。福布斯荣誉榜上的数据和基金名称每年都会发生变化。

表 3-20 投资组合换手率：2002 年福布斯荣誉榜上的共同基金

基金名称	年度投资组合 换手率 (%)	基金名称	年度投资组合 换手率 (%)
FPA Capital Fund (FPPTX)	23	Scudder-Dreman High Return Equity-A (KDHAX)	29
Mairs & Power Growth Fund (MPCFX)	8	Clipper Fund (CFIMX)	23
Muhlenkamp Fund (MUHLX)	11	Weitz Value Fund (WVALX)	13
Longleaf Partners Fund (LLPFX)	18	Third Avenue Value Fund (TAVFX)	16
Heartland Value Fund (HRTVX)	56	Dodge & Cox Stock Fund (DODGX)	10

资料来源：福布斯（Forbes 2003）。

基于表 3-20 中的数据，回答下述问题：

- (1) 利用 2002 年福布斯荣誉榜中的 10 大基金数据，计算区间投资组合换手率的总体均值。
- (2) 计算投资组合换手率的总体方差和总体标准差。
- (3) 解释本例中总体公式的使用原因。

解：

(1) $\mu = \frac{23 + 8 + 11 + 18 + 56 + 29 + 23 + 13 + 16 + 10}{10} = \frac{207}{10} = 20.7\%$ 。

(2) 在 $\mu = 20.7$ 的基础上，我们可以计算 $\sigma^2 = \frac{\sum_{i=1}^N (X_i - \mu)^2}{N}$ 。我们先计算表达式中的分子，再将其除以 $N = 10$ 。分子（偏离均值偏差的平方和）为

$$\begin{aligned} & (23 - 20.7)^2 + (8 - 20.7)^2 + (11 - 20.7)^2 + (18 - 20.7)^2 + \\ & (56 - 20.7)^2 + (29 - 20.7)^2 + (23 - 20.7)^2 + (13 - 20.7)^2 + \\ & (16 - 20.7)^2 + (10 - 20.7)^2 = 1784.1 \end{aligned}$$

因此， $\sigma^2 = \frac{1784.1}{10} = 178.41$ 。

为了计算标准差， $\sigma = \sqrt{178.41} = 13.357\%$ 。（方差的单位是百分比的平方，因此标准差的单位是百分比。）

① 《福布斯》杂志每年会选择出满足其荣誉榜特定标准的美国股票型共同基金。该标准与资本保值（在熊市中的表现）、管理持续性（基金必须聘用一个基金经理不少于 6 年），投资分散化程度，可投资性（不满足该要求的基金是指已经对于新投资者关闭的基金）和税后长期业绩表现有关。

(3) 如果总体明确地定义为在特定一年(2002)福布斯荣誉榜中的基金,且投资组合换手率是该特定一年区间中由福布斯所给出的数据,那么对于方差和标准差应用总体公式就是合理的。我们所得的结果 178.41 和 13.357 分别是 2002 年福布斯荣誉榜中基金横截面的年度投资组合换手率的方差和标准差。[⊖]

3.7.4 样本方差和样本标准差

3.7.4.1 样本方差

在投资管理的许多例子中,总体的一个子集或者样本是我们所能观测到的全部。当我们处理样本数据时,汇总性度量指标被称为统计量。用来度量样本数据离散度的统计量被称为样本方差。

- **样本方差计算公式 (sample variance formula)。** 样本方差 (sample variance) 为

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1} \quad (3-13)$$

式中, $\bar{X}$ 表示样本均值, n 表示样本中观测值的数目。

式(3-13)告诉我们通过下述步骤来求解样本方差:

1. 计算样本均值 $\bar{X}$ 。
2. 计算每个观测值偏离样本均值的平方偏差值 $(X_i - \bar{X})^2$ 。
3. 将所有偏离均值的偏差值平方加总: $\sum_{i=1}^n (X_i - \bar{X})^2$ 。
4. 将偏离均值的偏差值平方和除以 $n - 1$: $\sum_{i=1}^n (X_i - \bar{X})^2 / (n - 1)$ 。

我们将在例 3-12 中举例说明样本方差和样本标准差的计算方法。

我们用 s^2 来表示样本方差,以区分总体方差 σ^2 。样本方差的计算公式除了使用样本均值 $\bar{X}$ 代替总体均值 μ 以及除数不同之外,其他都与总体方差的计算公式相类似。在总体方差的例子中,我们所除的是总体的规模 N 。然而,对于样本方差,我们所除的是样本规模减 1,或者是 $n - 1$ 。通过使用 $n - 1$ (而不是 n) 作为除数,我们改善了样本方差的统计性质。用统计术语来说,式(3-13)所定义的样本方差是总体方差的无偏统计量。[⊖] $n - 1$ 这个数值被也认为是估计总体方差的自由度。为了用 s^2 来估计总体方差,我们必须首先算出样本均值。一旦我们计算出样本均值,那么只有 $n - 1$ 个独立的偏离均值的偏差值。

3.7.4.2 样本标准差

和我们计算总体标准差一样,我们可以通过求解样本方差的正的平方根来计算样本标准差。

- **样本标准差的计算公式 (sample standard deviation formula)。** 样本标准差 (sample standard deviation) s 为

⊖ 事实上,我们不能用荣誉榜中的基金数据来恰当地估计任何不同定义的投资组合换手率的总体方差(作为一个例子),因为荣誉榜中的基金数据不是从任何更大的美国股票型共同基金总体中所取出的随机样本。

⊖ 我们将在抽样一章中对于这一概念进行深入讨论。

$$s = \sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1}} \tag{3-14}$$

式中， $\bar{X}$ 表示样本均值， n 表示样本中观测值的数目。

为了计算样本标准差，我们首先根据给定的步骤计算样本方差，然后我们再求样本方差的平方根。例3-12举例说明了我们之前介绍的两家共同基金样本方差和样本标准差的计算方法。

例 3-12 样本方差和样本标准差的计算

在计算出例3-7中两个共同基金的几何平均和算术平均收益率之后，我们在例3-10中计算出了度量这两家基金离散度的两个指标——收益率的极差和平均绝对偏差。我们现在计算这两家基金收益率的样本方差和样本标准差。

表 3-15（续） 两家共同基金的总收益率，1998~2002 年

年份	所选的美国股票 (SLASX) (%)	T. Rowe 价格股票收益 (PRFDX) (%)	年份	所选的美国股票 (SLASX) (%)	T. Rowe 价格股票收益 (PRFDX) (%)
1998	16.2	9.2	2001	-11.1	1.6
1999	20.3	3.8	2002	-17.0	-13.0
2000	9.3	13.1			

资料来源：AAIL。

基于上面表3-15所给出的数据，回答下列问题：

- (1) 计算 (A) SLASX 和 (B) PRFDX 收益率的样本方差。
- (2) 计算 (A) SLASX 和 (B) PRFDX 收益率的样本标准差。
- (3) 比较两家基金中每个基金收益率的标准差和平均绝对偏差这两个不同收益率离散度的度量指标间的差别。

解：

(1) 为了计算样本方差，我们利用式 (3-13)。(方差的结果都表示成百分比的平方。)

A. SLASX

- 1) 样本均值为 $\bar{R} = \frac{16.2 + 20.3 + 9.3 - 11.1 - 17.0}{5} = \frac{17.7}{5} = 3.54\%$ 。
- 2) 偏离均值偏差的平方分别为：
 $(16.2 - 3.54)^2 = (12.66)^2 = 160.2756$
 $(20.3 - 3.54)^2 = (16.76)^2 = 280.8976$
 $(9.3 - 3.54)^2 = (5.76)^2 = 33.1776$
 $(-11.1 - 3.54)^2 = (-14.64)^2 = 214.3296$
 $(-17.0 - 3.54)^2 = (-20.54)^2 = 421.8916$
- 3) 偏离均值偏差的平方和为 $160.2756 + 280.8976 + 33.1776 + 214.3296 + 421.8916 = 1110.4720$ 。
- 4) 将偏离均值偏差平方和除以 $n - 1$ 为： $1110.5720 / (5 - 1) = 1110.5720 / 4 = 277.6430$ 。

B. PRFDX

- 1) 样本均值为 $\bar{R} = \frac{9.2 + 3.8 + 13.1 + 1.6 - 13.0}{5} = \frac{14.7}{5} = 2.94\%$ 。
- 2) 偏离均值偏差的平方分别为：
- $(9.2 - 2.94)^2 = (6.26)^2 = 39.1876$
- $(3.8 - 2.94)^2 = (0.86)^2 = 0.7396$
- $(13.1 - 2.94)^2 = (10.16)^2 = 103.2256$
- $(1.6 - 2.94)^2 = (-1.34)^2 = 1.7956$
- $(-13.0 - 2.94)^2 = (-15.94)^2 = 254.0836$
- 3) 偏离均值偏差的平方和为 $39.1876 + 0.7396 + 103.2256 + 1.7956 + 254.0836 = 399.032$ 。
- 4) 将偏离均值偏差平方和除以 $n - 1$ 为： $399.032 / (5 - 1) = 399.032 / 4 = 99.758$ 。
- (2) 为了求解标准差，我们计算方差的正的平方根。
- A. 对于 SLASX， $\sigma = \sqrt{277.6430} = 16.66\%$ 或者 16.7% 。
- B. 对于 PRFDX， $\sigma = \sqrt{99.758} = 9.99\%$ 或者 10.0% 。
- (3) 表 3-21 汇总了第 2 题标准差和例 3-10 所得到的 MAD 的结果。

表 3-21 两家共同基金：标准差和平均绝对偏差间的比较

基金	标准差 (%)	平均绝对偏差 (%)
SLASX	16.7	14.1
PRFDX	10.0	6.9

注意到平均绝对偏差小于标准差。平均绝对偏差总是小于等于标准差，因为标准差相对于较小的偏差，给予较大的偏差以更大的权重。（记住，方差是平方的。）

因为标准差是度量偏离算术平均数离散度的指标，我们在汇总描述数据时通常将算术平均数和标准差一起报告给出。当我们处理表示百分比变化的时间序列数据时，给出几何平均数——表示复利增长率，也是对我们非常有帮助的。表 3-22 给出了不同股票收益率序列的历史几何平均收益率和算术平均数，以及历史收益率的标准差。由于我们的这些统计量以名义收益率（而不是通胀调整后的收益率）表述的，因此我们可以观测到原始收益率的大小。

表 3-22 股票市场收益率：均值和标准差

收益率序列	几何平均数 (%)	算术平均数 (%)	标准差 (%)
伊博森协会序列：1926 ~ 2002 年			
S&P 500 (年度)	10.20	12.20	20.49
S&P 500 (月度)	0.81	0.97	5.65
蒂姆森等 (2002) 序列 (年度)：1900 ~ 2000 年			
澳大利亚	11.9	13.3	18.2
比利时	8.2	10.5	24.1
加拿大	9.7	11.0	16.6
丹麦	8.9	10.7	21.7
法国	12.1	14.5	24.6

(续)

收益率序列	几何平均数 (%)	算术平均数 (%)	标准差 (%)
德国	9.7	15.2	36.4
爱尔兰	9.5	11.5	22.8
意大利	12.0	16.1	34.2
日本	12.5	15.9	29.5
荷兰	9.0	11.0	22.7
南非	12.0	14.2	23.7
西班牙	10.0	12.1	22.8
瑞典	11.6	13.9	23.5
瑞士	7.6	9.3	19.7
英国	10.1	11.9	21.8
美国	10.1	12.0	19.9

资料来源：伊博森 EnCorrAnalyzer；蒂姆森等

3.7.5 半方差、半离差及其相关概念

一项资产收益率的方差或者标准差经常被解释成资产风险的度量指标。收益率的方差和标准差考虑了均值上方和下方的收益率，但是投资者常常只关心下方风险，即低于均值的收益率。结果，分析师们提出了半方差、半离差和相关的关注下方风险的度量离散度的指标。半方差 (semi-variance) 定义为低于均值以下部分平均平方偏差。半离差 (semideviation) (有时也被称为半标准差) 是半方差的正的平方根。为了计算半方差，我们采取如下步骤：

- 1. 计算样本均值。
- 2. 识别出小于均值的观测值 (去除大于等于均值的观测值)；假设有 n^* 个观测值小于均值。
- 3. 计算偏离均值偏差为负的偏差的平方和 (使用 n^* 个小于均值的观测值)。
- 4. 将第 3 步所求得的负偏差的平方和除以 $n^* - 1$ 。半方差的公式为：

$$\sum_{\text{对于所有 } X_i < \bar{X}} (X_i - \bar{X})^2 / (n^* - 1)$$

利用所选美国股票 (收益率 (以百分比计) 分别为 16.2, 20.3, 9.3, -11.1 和 -17.0) 的例子，我们之前计算出均值收益率为 3.54%。两个收益率 -11.1 和 -17.0 小于 3.54 ($n^* = 2$)。我们计算偏离均值负偏差的平方和为 $(-11.1 - 3.54)^2 + (-17.0 - 3.54)^2 = (-14.64)^2 + (-20.54)^2 = 214.3296 + 421.8916 = 636.2212$ 。由于 $n^* - 1 = 1$ ，我们得到半方差为 $636.2212 / 1 = 636.2212$ ，于是半离差约为 $\sqrt{636.2212} = 25.2\%$ 。半离差 25.2% 比 16.7% 的标准差要大。因此，从下行风险的角度来看，标准差低估了风险。

在实践中，我们可能关心低于除均值以外其他水平的收益率的值。例如，如果我们的收益率目标为每年 10%，我们可能特别关心低于每年 10% 的收益率。我们可以将 10% 称为目标值。目标半方差 (target semivariance) 的名称是指低于给定目标值的偏差平方平均，而目标半离差 (target semideviation) 是其正的平方根。为了计算样本目标半方差，我们先在第 1 步确定一个目标值。在识别出低于目标值的观测值之后，我们求解低于目标值的负偏差的平方和并将其除以低于目标值观测值数目减 1 的值。目标半方差的计算公式为：

$$\sum_{\text{对于所有 } X_i < B} (X_i - B)^2 / (n^* - 1)$$

式中， B 是目标值，并且 n^* 是低于目标值观测值的数目。当目标值为 10% 时，我们在所选美国股票的例子中找到 3 个低于目标值的收益率（9.3，-11.1 和 -17%）。故目标半方差为 $[(9.3 - 10.0)^2 + (-11.1 - 10.0)^2 + (-17.0 - 10.0)^2] / (3-1) = 587.35$ ，且目标半离差约为 $\sqrt{587.35} = 24.24\%$ 。

当收益率分布是对称时，半方差是方差的一个常数比例（一半），故这两个度量指标是同样有效的。对于不对称的分布，方差和半方差将给出不同的潜在风险排序。[Ⓐ]半方差（或半离差）和目标半方差（或目标半离差）虽然贴近直观，但是在数学上它们比方差难以处理。[Ⓑ]方差或者标准差出现在许多最常用的金融风险的概念定义中，如夏普比率和贝塔。可能由于这些原因，使得方差（或标准差）在投资实践中的应用更加地频繁。

3.7.6 切比雪夫不等式

俄罗斯数学家切比雪夫（Pafnuty Chebyshev）利用标准差作为离散度的度量指标，推导出了一个不等式。该不等式给出了在均值周围 k 个标准差之内数值所占的比例。

- **切比雪夫不等式的定义。**根据切比雪夫不等式，在算术平均数 k 个标准差内观测值的比例至少为 $1 - 1/k^2$ （对于所有 $k > 1$ ）。

表 3-23 给出了落入样本均值给定数目标标准差的范围内的观测值所占的比例。

表 3-23 由切比雪夫不等式所得到的比例

k	包含样本均值的区间	比例 (%)	k	包含样本均值的区间	比例 (%)
1.25	$\bar{X} \pm 1.25s$	36	2.50	$\bar{X} \pm 2.50s$	84
1.50	$\bar{X} \pm 1.50s$	56	3.00	$\bar{X} \pm 3s$	89
2.00	$\bar{X} \pm 2s$	75	4.00	$\bar{X} \pm 4s$	94

注：标准差用符号 s 表示。

例如，当 $k = 1.25$ 时，该不等式表示位于 $\pm 1.25s$ 中的观测值的最小比例为 $1 - 1/(1.25)^2 = 1 - 0.64 = 0.36$ 或者 36%。

由切比雪夫不等式得到的最频繁引用的事实是无论数据服从什么分布，包含均值两倍标准差区间包含不少于 75% 的观测值，而包含均值两倍标准差区间包含不少于 89% 的观测值。

切比雪夫不等式的重要性在于其一般性。无论分布的形状，对于样本数据还是总体数据，对于离散数据还是连续数据，该不等式都成立。正如我们将在抽样一章中所看到的，如果我们假设所抽取样本的总体服从一个特定的称为正态分布的分布，那么我们可以得到非常精确的区间清单。然而，通常我们无法明确地假设数据的分布情况。

下面的例子举例说明了切比雪夫不等式的应用。

Ⓐ 对于负偏的收益率分布，半方差比方差的一半要大；对于正偏的收益率分布，半方差比方差的一半要小。参见 Estrada（2003）。我们将在本章的后面讨论偏度。
Ⓑ 如概率概念一章和投资组合概念一章中所讨论的，我们能求出一个投资组合的方差。该方差是各个证券的方差和投资组合中各个证券间相关关系的直接函数。半方差和目标半方差的求解过程并不相同。我们也不能对半方差或目标半方差求倒数。

例 3-13 切比雪夫不等式的应用

根据表 3-22，在 1926 ~ 2002 年间标准普尔 500 月收益率（共 924 个月度观测值）的算术平均数和标准差分别为 0.97% 和 5.65%。利用这些信息，回答下述问题：

- (1) 计算根据切比雪夫不等式得到的包含至少 75% 的月收益率的区间端点。
- (2) 第 1 题中，根据切比雪夫不等式所计算得到的区间内，观测值的最小值和最大值分别是多少？

解：

(1) 根据切比雪夫不等式，至少 75% 的观测值必位于均值的两个标准差之内， $\bar{X} \pm 2s$ 。对于标准普尔 500 月收益率序列，我们有 $0.97\% \pm 2(5.65\%) = 0.97\% \pm 11.30\%$ 。因此，包含至少 75% 观测值的区间下端点为 $0.97\% - 11.30\% = -10.33\%$ ，而上端点为 $0.97\% + 11.30\% = 12.27\%$ 。

(2) 对于大小为 924 的样本，至少有 $0.75(924) = 693$ 个观测值位于在第 1 题中所计算出的 $-10.33\% \sim 12.27\%$ 的区间之中。切比雪夫不等式给出了观测值必须落在包含均值给定区间的最小比例，但是它没有给出最大比例。表 3-4 给出了标准普尔 500 月收益率的频数分布，具体摘录如下。摘取的表中的数据与切比雪夫不等式的预测是一致的。从 $-10\% \sim 12.0\%$ 的一组区间宽度仅稍稍比两倍的标准差区间 $-10.33\% \sim 12.27\%$ 窄一些。共有 886 个观测值（约为 96% 的观测值）落在区间 $-10.0\% \sim 12.0\%$ 中。

表 3-4 （摘取的）1926 年 1 月 ~ 2002 年 12 月标准普尔 500 月度总收益率的频数分布

收益率区间 (%)	绝对频数	收益率区间 (%)	绝对频数
-10.0 ~ -8.0	20	2.0 ~ 4.0	153
-8.0 ~ -6.0	30	4.0 ~ 6.0	126
-6.0 ~ -4.0	54	6.0 ~ 8.0	58
-4.0 ~ -2.0	90	8.0 ~ 10.0	21
-2.0 ~ 0.0	138	10.0 ~ 12.0	14
0.0 ~ 2.0	182		886

3.7.7 变异系数

我们之前提到标准差要比方差更加易于解释，这是因为标准差使用了与观测值相同的度量单位。然而我们可能有时会发现无法解释不同组数据间相对变动程度的标准差的含义，这或者是因为数据集有不同的均值，或者是因为数据集以不同的度量单位进行计量。在这一节中我们将介绍相对离散度的度量指标，变异系数在该问题中非常有用。相对离散度（relative dispersion）是相对于一个参考点或者基准的偏离大小。

我们可以用两个假设的公司样本的例子来说明拥有不同均值数据集的标准差解释问题。第 1 个样本由小公司组成，其包括在 2003 年销售额为 50 000 000 欧元、75 000 000 欧元、65 000 000 欧元和 90 000 000 欧元的公司。第 2 个样本由大公司组成，其包括在 2003 年销售额为 800 000 000 欧元、825 000 000 欧元、815 000 000 欧元和 840 000 000 欧元的公司。我们可以利用式 (3-14) 证实：两个样本的销售额的标准差均为 16 800 000 欧元。^①在第 1 个样本中，最大的观测值为 90 000 000

① 第 2 个样本是由第 1 个样本中每个观测值加上 750 000 000 欧元而得到的。当我们将一个常数加到每个观测值上时，该数据的标准差（和方差）具有保持不变的性质。

欧元，它比最小的观测值 50 000 000 欧元大 80%。第 2 个样本中，最大的观测值仅仅比最小的观测值大 5%。简单地来看，16 800 000 欧元的标准差对于第 1 个样本（2003 年销售均值 75 000 000 欧元）来说，是一个较大的变动程度。但是，其相对于第 2 个样本（2003 年销售均值 820 000 000 欧元）来说，却是一个较少的变动程度。

变异系数在刚才描述的这种情况下就对我们很有帮助。

- **变异系数的计算公式**（coefficient of variation formula）。变异系数（coefficient of variation, CV）是一组观测值的标准差与其均值的比值：[⊖]

$$CV = s/\bar{X} \tag{3-15}$$

式中， s 表示样本标准差， $\bar{X}$ 表示样本均值。

当观测值是收益率时，变异系数度量的是每单位平均收益率所承担的风险（标准差）。变异系数是以观测值变动大小相对于它们均值的比率来表示的，它使得我们可以直接在不同数据间进行离散度的比较。为了对于量纲进行纠正，变异系数是一个无量纲的度量指标（即它没有度量单位）。

我们可以用我们之前两个公司样本的例子来对于变异系数的应用进行说明。第 1 个样本的变异系数为 $(\text{€} 16\,800\,000)/(\text{€} 70\,000\,000) = 0.24$ ；而第 2 个样本的变异系数为 $(\text{€} 16\,800\,000)/(\text{€} 80\,000\,000) = 0.02$ 。这一结果证实了我们的直觉，即第 1 个样本比第 2 个样本销售额的变动更大。注意 0.24 和 0.02 是纯数，因为它们都是没有度量单位的（我们用的是标准差除以和标准差同样单位的均值）。如果我们需要比较以不同度量单位数据集之间的离散度，那么变异系数可能是有用的，因为它不含有度量单位。例 3-14 演示了变异系数的计算过程。

例 3-14 变异系数的计算

表 3-24 汇总了多个主要美国资产类的平均年收益率和标准差，其中利用伊博森 EnCorr Analyzer 中的一个选项将月度收益率统计量转换成了年度的相应指标。

表 3-24 美国资产类算术平均年收益率及其标准差（1926 ~ 2002 年）

资产类	算术平均 收益率 (%)	收益率的 标准差 (%)	资产类	算术平均 收益率 (%)	收益率的 标准差 (%)
标准普尔 500	12.3	21.9	美国长期政府债	5.8	8.2
美国小盘股	16.9	35.1	美国 30 天国库券	3.8	0.9
美国长期公司债	6.1	7.2			

资料来源：伊博森 EnCorr Analyzer。

利用表 3-24 中的信息，回答下列问题：

- (1) 计算每个给定资产类的变异系数。
- (2) 利用 CV 作为相对离散度的度量指标，对资产类从风险最大到最小进行排序。
- (3) 确定标准普尔 500 和美国小盘股的绝对和相对风险间是否存在较大的差别。利用标准差作为绝对风险的度量指标，而用 CV 作为相对风险的度量指标。

⊖ 读者也将会遇到把 CV 定义成 $100 (S/\bar{X})$ 的情况，这里 CV 是以百分比表示的。

解:

(1) 标准普尔 500: $CV = 21.9\% / 12.3\% = 1.780$

美国小盘股: $CV = 35.1\% / 16.9\% = 2.077$

美国长期公司债: $CV = 7.2\% / 6.1\% = 1.180$

美国长期政府债: $CV = 8.2\% / 5.8\% = 1.414$

美国 30 天国库券: $CV = 0.9\% / 3.8\% = 0.237$

(2) 基于变异系数 (CV), 具体的排序为: 美国小盘股 (风险最大), 标准普尔 500, 美国长期政府债, 美国长期公司债和美国 30 天国库券 (风险最小)。

(3) 如标准差和变异系数所度量的, 美国小盘股要比标准普尔 500 的风险更大。然而, 变异系数揭示出小盘股和标准普尔 500 收益率变动的差别要小于仅由标准差所得到的两者差别。小盘股的标准差比标准普尔 500 大 $(35.1 - 21.9) / 21.9 = 0.603$ 或者 60% 左右, 而相比而言, 变异系数的差别仅为 $(2.077 - 1.780) / 1.780 = 0.167$ 或者 17%。

3.7.8 夏普比率

虽然变异系数被设计为相对离散度的度量指标, 但是它的倒数却反映了单位风险的收益率情况, 因为收益率的标准差经常被用来作为投资风险的度量指标。例如, 一个平均月收益率为 1.19%, 标准差为 4.42% 的投资组合, 其变异系数的倒数为 $1.19\% / 4.42\% = 0.27$ 。这一结果表明了每单位的标准差带来了 0.27% 的收益率。

一个更加精确的收益—风险的度量指标考虑了一种无风险收益率的存在。所谓的无风险利率是指其标准差几乎为 0 的一种收益率。在存在无风险资产的情况下, 一个投资者可以选择风险投资组合 p , 并将其与无风险资产混合以达到任何以标准差 s_p 作为度量的要求风险水平。考虑一个将平均收益率作为纵轴, 而将标准差作为横轴的图像。任何投资组合 p 和无风险资产的组合将位于一条射线上, 其斜率等于 (平均收益率 - 无风险收益率) 除以 s_p 。这条射线给出了投资者所能选择的资产中单位风险具有最大收益 (超过无风险收益率的收益) 的选择, 因此, 其具有最高的斜率。对于一个投资组合 p 的超额收益率与收益率标准差的比率 (经过 p 直线的斜率) 是用来度量投资组合业绩的一个数量指标, 它被称为夏普比率 (以表明其发明者为威廉·夏普)。

- 夏普比率的计算公式 (Sharpe ratio formula)。基于历史收益率, 一个投资组合 p 的夏普比率被定义为:

$$S_h = \frac{\bar{R}_p - \bar{R}_f}{s_p} \quad (3-16)$$

式中, $\bar{R}_p$ 表示投资组合的平均收益率, $\bar{R}_f$ 表示无风险资产的平均收益率, 而 s_p 表示该投资组合收益率的标准差。[⊖]

夏普测度的分子是样本区间上投资组合平均收益率减去无风险资产平均收益率的值。 $\bar{R}_p - \bar{R}_f$

⊖ 该等式表示的是事后的或者是历史的夏普比率。我们也可以基于对平均收益率、无风险收益率和收益率标准差的预测来考虑一个投资组合未来的夏普比率。我们也可能遇到夏普比率的另一种计算方法, 即用 (投资组合收益率 - 无风险资产收益率) 序列的标准差而不是投资组合收益率序列的标准差作为其分母的计算方法; 在实践中, 这两种标准差的计算方法一般得到的结果很接近。想要了解更多关于夏普比率 (它也被称为夏普测度, 收益—波动比率和超额收益—波动比率) 的信息, 参见 Elton, Gruber, Brown 和 Goetzmann (2003) 和 Sharpe (1994)。

项度量了投资者承担额外风险所能获得的额外报酬。我们将这个差称为投资组合 p 的平均超额收益率 (mean excess return)。因此, 夏普比率以每单位风险 (以收益率标准差来度量的) 的平均超额收益率来度量收益。那些风险厌恶的投资者在作决策时, 只会偏好那些使得投资组合的夏普比率较大的投资组合的平均收益率和标准差, 而不会去选择那些使得夏普比率较小的投资组合的平均收益率和标准差。

为了举例说明夏普比率的计算方法, 我们考虑之前在表 3-24 中给出的 1926 ~ 2002 年间标准普尔 500 和美国小盘股的业绩表现的例子。将美国国库券的平均收益率作为无风险收益率的代表, 我们可以求得:

$$\text{标准普尔 500: } S_h = \frac{12.3 - 3.8}{21.9} = 0.39$$

$$\text{美国小盘股: } S_h = \frac{16.9 - 3.8}{35.1} = 0.37$$

虽然美国小盘股的平均收益率较大, 但是从夏普比率的情况来看, 它们的表现却稍弱于标准普尔 500。

夏普比率是业绩评估中的主要指标。在此, 我们必须说明两个在其应用中所需要注意的问题。一个和负的夏普比率的解释有关, 另一个与该概念的局限性有关。

金融理论告诉我们, 从长期来看, 至少在风险投资组合完全分散化的情况下, 投资者所承担的额外风险应该由超过无风险收益率的额外平均收益率来进行补偿。如果投资者确实获得了相应的补偿, 那么夏普比率的分子就应该为正。然而, 当夏普比率的计算期为熊市占主导的多个时间段时, 我们经常发现所得的投资组合的夏普比率为负。当出现负的夏普比率时, 这就要引起我们的注意。当夏普比率为正时, 在其他情况不变的情况下, 一个投资组合的夏普比率将会随着我们承担风险的增加而减少。这一结果对于一个经过风险调整后的业绩度量指标来说是非常直观的。然而, 当夏普比率为负时, 风险的增加却会造成数值上更大的夏普比率 (例如, 风险加倍将会使得夏普比率从 -1 增长到 -0.5)。因此, 在对于拥有负的夏普比率的投资组合进行比较时, 我们一般不能将大的夏普比率 (更接近于 0 的夏普比率) 解释为对我们更有利的风险调整后的业绩情况。^① 实践中, 为了使得在这种情况下能够使用夏普比率进行解释与比较, 我们可能需要增加评估区间段以使一个或者更多的夏普比率为正; 我们也可能要考虑一个不同的业绩评估度量标准。

夏普比率这个概念的局限性在于它仅仅考虑了风险的一个方面, 收益率的标准差。标准差作为拥有近似对称收益率分布的投资组合策略的风险度量指标是最适合的。而带有选择元素的策略会具有不对称的收益率分布。与此相关的, 一个投资策略可能会带来经常性的小收益和不经常性极大损失的可能性。^② 这类策略有时被描述成在一个推土机前捡硬币 (你能保证长期的较小收益, 但是推土机的一个突然启动就会把你碾得粉碎。——译者注); 举个例子来说, 有些对冲基金策略就倾向于产生这种收益率的分布情况。如果在该策略起作用的区间 (大的损失没有发生) 内进行计算, 则该类策略会有一个较高的夏普比率。在这个例子中, 夏普比率可能给出一个过于乐观的经过风险调整后的业绩情况, 因为标准差不能完全度量给定的风险。^③ 因此, 在应用夏普比率对基金经理进行评估之前, 我们应该先判断一下标准差是否能够充分地反映该经理的投资策略所承担的风险。

例 3-15 举例说明了夏普比率在投资组合业绩评价中如何进行计算。

① 然而, 如果标准差是相等的, 那么具有负的且更接近 0 的夏普比率的投资组合对我们更有利。

② 这一表述描述的是偏度为负的收益率分布。我们将在本章的后面讨论偏度的概念。

③ 需要了解更多信息, 参见 Amin 和 Kat (2003)

例 3-15 夏普比率的计算

在之前的例子中，我们计算了两家共同基金（选择美国股票（SLASX）和普信股票收益基金（PRFDX））截至 2002 年 12 月为止 5 年的不同的统计量。表 3-25 汇总了更长区间（截至 2002 年 10 月）的对于两家共同基金所选择的统计量。

表 3-25 共同基金平均收益率及其标准差

基金	算术平均数 (%)	收益率的标准差 (%)
SLASX	12.58	19.44
PRFDX	11.64	13.65

资料来源：AAIL

美国 30 天国库券利率经常作为无风险利率的代理变量。表 3-26 给出了 1993 ~ 2002 年间国库券的年收益率。

表 3-26 年化的美国 30 天国库券收益率（1993 ~ 2002 年）

年份	收益率 (%)	年份	收益率 (%)
1993	2.90	1998	4.86
1994	3.90	1999	4.68
1995	5.60	2000	5.89
1996	5.21	2001	3.83
1997	5.26	2002	1.65

资料来源：伊博森协会。

利用表 3-25 和表 3-26 中的信息，回答下述问题：

- (1) 计算 1993 ~ 2002 年间 SLASX 和 PRFDX 的夏普比率。
- (2) 说明根据夏普比率这个度量指标，哪一个基金在这个区间具有更好的风险调整后业绩。

解：

(1) 我们已经有了投资组合平均收益率及其标准差。而 1993 ~ 2002 年的平均无风险年收益率（用美国国库券作为代理变量）为 $(2.90 + 3.90 + 5.60 + 5.21 + 5.26 + 4.86 + 4.68 + 5.89 + 3.83 + 1.65) / 10 = 43.78 / 10 = 4.38\%$

$$\text{SLASX: } S_{h,\text{SLASX}} = \frac{12.58 - 4.38}{19.44} = 0.42$$

$$\text{PRFDX: } S_{h,\text{PRFDX}} = \frac{11.64 - 4.38}{13.65} = 0.53$$

(2) PRFDX 比 SLASX 在该区间内有更高的正的夏普比率。故根据夏普比率这个指标，PRFDX 的业绩更好。

3.8 收益率分布的对称性和偏度

均值和方差可能并不能充分描述投资收益率的分布情况。例如，在计算方差时，由于对偏离均值偏差进行了平方，因此，我们不知道大的偏差更可能为正还是为负。我们需要有除度量集中趋势和离散度的其他指标来揭示分布的其他重要特征。其中，一个分析师感兴趣的重要指标就是收益率分布的对称程度。

如果一个收益率分布是关于其均值对称的，那么分布的每一侧都是另一侧的镜面映像。因此，损失和收益数值相同的区间将表现出相同的频数。例如，位于 -5% ~ -3% 区间内损失的频

数与位于3%~5%区间内收益的频数相同。

一个最重要的分布为正态分布，其具体表现为图3-6。这个对称的钟形分布在投资组合选择的均值方差模型中发挥着重要的作用，它也被广泛应用于金融风险管理之中。正态分布具有以下这些特点：

- 它的均值和中位数是相等的。
- 它可以完全由两个参数（其均值和方差）来描述。
- 大约68%的观测值位于其均值的正负一个标准差之间；95%的观测值位于其均值的正负两个标准差之间；99%的观测值位于其均值的正负3个标准差之间。

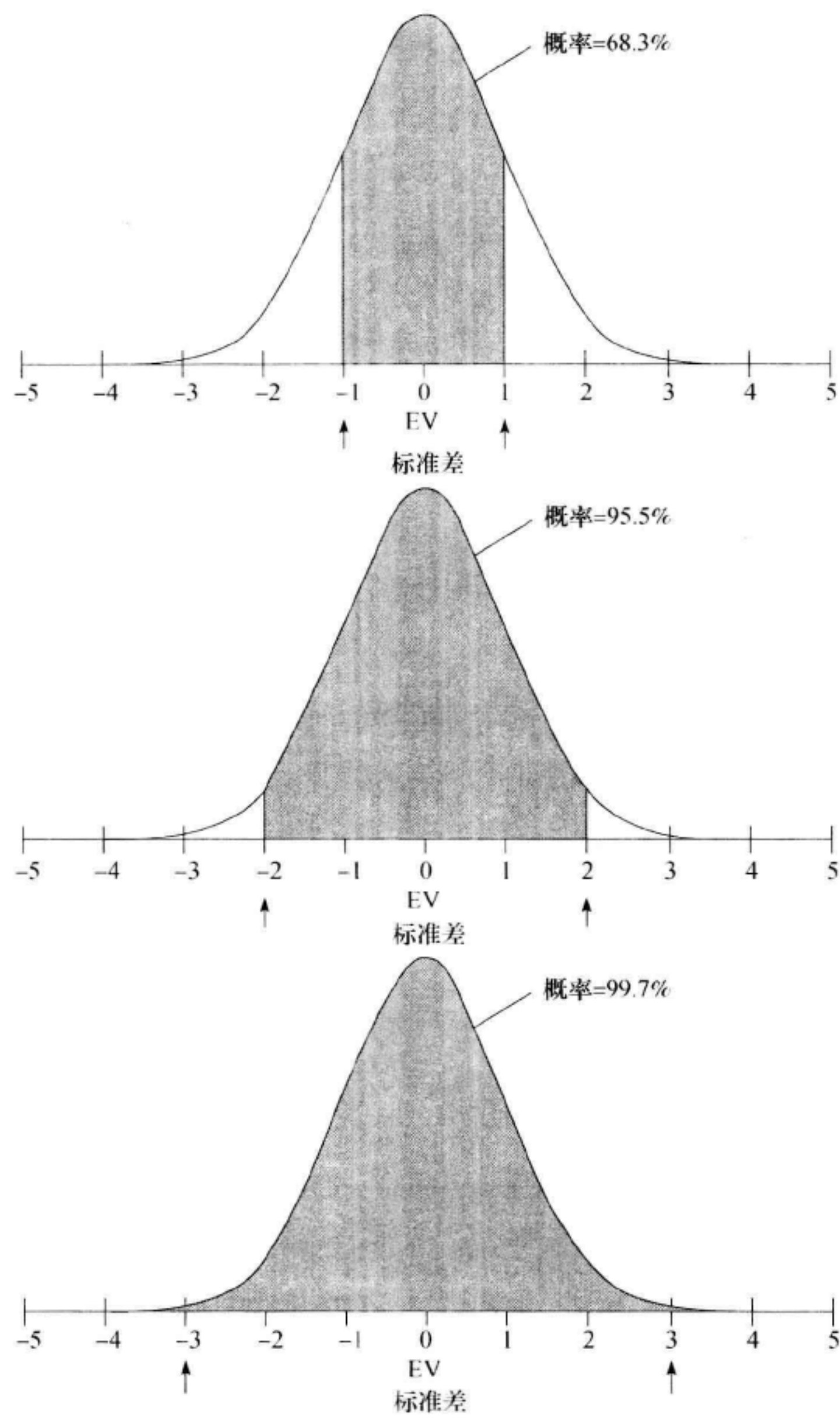


图 3-6 正态分布的性质（EV = 期望值）

资料来源：Reprinted with permission from *Fixed Income Readings for the Chartered Financial Analyst® Program*. Copyright 2000, Frank J. Fabozzi Associates, New Hope, PA.

一个分布如果不是对称的，则被称为偏的 (skewed)。一个正偏的收益率分布会较频繁地出现小损失和一些极大的收益。一个负偏的收益率分布会较频繁地出现小收益和一些极大的损失。图 3-7 展现了正偏和负偏的分布情况。正偏的分布表现出右侧长尾；而负偏的分布具有左侧长尾。对于正偏的单峰分布，其众数要小于中位数，而中位数又要小于均值。而对于负偏的单峰分布，其均值要小于中位数，而中位数又要小于众数。^① 投资者应该更喜欢正偏 (分布)，因为其平均收益率高于中位数。相对于平均收益率，正偏具有一个有限的 (虽然是经常发生的) 损失，而具有一个无限的 (但是，它不是经常发生的) 收益。

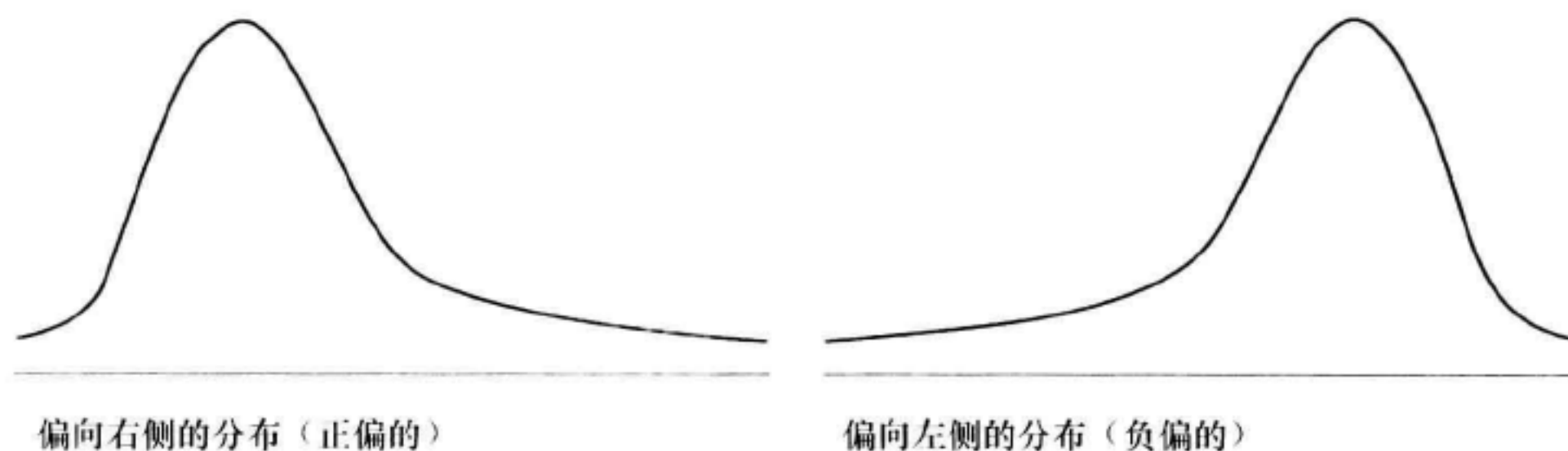


图 3-7 有偏分布的性质

资料来源: Reprinted with permission from *Fixed Income Readings for the Chartered Financial Analyst® Program*. Copyright 2000, Frank J. Fabozzi Associates, New Hope, PA.

偏度是对度量偏斜程度统计指标所给予的名称。(“偏度”这个词有时也交替地用来表示“偏斜程度”。) 与方差相同，偏度的计算也使用每个观测值偏离均值的偏差值。偏度 (skewness) (有时也被称为相对偏度) 被计算为偏离均值的偏差值的立方平均值，再除以标准差的立方来进行标准化，从而使得该度量指标不受量纲的影响。^② 根据该度量指标的定义，一个对称分布的偏度为 0，一个正偏的分布具有正的偏度，而一个负偏的分布具有负的偏度。

我们可以通过关注该指标的分子来说明该指标背后的原理。三次方不同于平方，它保持了偏离均值偏差值的符号。如果一个分布是正偏的 (其均值大于中位数)，那么超过一半的偏离均值的偏差为负，而为正的偏差值少于一半。为了使得其偏差立方和为正，其损失必须很小且经常发生，而其收益应该不经常发生但是数值很大。因此，如果偏度为正的话，正偏差的平均大小一定比负偏差的平均大小要大。

一个简单的例子可以说明一个对称分布的偏度 (度量指标) 等于 0。假设我们有如下这些数据：1, 2, 3, 4, 5, 6, 7, 8 和 9。均值的结果为 5，而偏差值为 -4, -3, -2, -1, 0, 1, 2, 3 和 4。偏差的立方为 -64, -27, -8, -1, 0, 1, 8, 27 和 64，可见其总和为 0。因此，偏度的分子等于 0 (故偏度其本身也为 0)。这证实了我们之前的观点。下面我们给出计算样本偏度的具体公式。

- **样本偏度计算公式 (sample skewness formula)**。样本偏度 (sample skewness) (也被称为样本相对偏度 (sample relative skewness)) S_k 为

$$S_k = \left[\frac{n}{(n-1)(n-2)} \right] \frac{\sum_{i=1}^n (X_i - \bar{X})^3}{s^3} \quad (3-17)$$

① 为了帮助记忆，在这个例子中，均值、中位数和众数以和它们在字典中排列顺序相同的顺序出现。

② 我们这里讨论的是矩偏度系数。一些教科书介绍的是皮尔逊偏度系数，其等于 3 (均值 - 中位数) / 标准差，它的缺点在于其涉及中位数的计算。

式中， n 表示样本中观测值的数目，而 s 表示样本标准差。[⊖]

式 (3-17) 的代数符号指出了偏斜的方向，一个负的 S_k 表示一个负偏的分布，而一个正的 S_k 表示一个正偏的分布。注意到当 n 变大时，该表达式将退化为平均立方偏差， $\dot{S}_k \approx \left(\frac{1}{n}\right)$

$\frac{\sum_{i=1}^n (X_i - \bar{X})^3}{s^3}$ 。作为一个参考，对于样本容量为 100 或者更大的来源于正态分布的样本来说， ± 0.5 的偏度系数会被认为出乎意料地大。

表 3-27 列出了标准普尔 500 年收益率和月收益率的一些汇总统计量。之前我们已经讨论过了平均收益率及其标准差，而我们将在不久之后简要地讨论一下峰度。

表 3-27 标准普尔 500 年度和月度的总收益率 (1926 ~ 2002 年)：汇总统计量

收益率序列	期数	算术平均数 (%)	标准差 (%)	偏度	超额峰度
S&P 500 (年度)	77	12.20	20.49	-0.294 3	-0.220 7
S&P 500 (月度)	924	0.97	5.65	0.396 4	9.464 5

资料来源：伊博森 EnCorrAnalyzer。

表 3-27 表明标准普尔 500 在该区间的年收益率是负偏的，而月收益率是正偏的，并且月收益率的偏度绝对值更大。我们将会发现其他市场序列的收益率分布形状常常是依赖于所考察的持有区间的。

一些研究者认为投资者应该在其他条件相同的情况下，更偏好正的偏度，即他们应该偏好拥有提供相对较频繁出现的不寻常高收益分布的投资组合。[⊖]不同投资策略可能倾向于提供不同种类和大小的收益率偏度。例 3-16 举例说明了一个管理的投资组合偏度的计算方法。

例 3-16 共同基金偏度的计算

表 3-28 给出了普信股票收益基金 (PRFDX) 10 年的年收益率。

表 3-28 年收益率：普信股票收益基金 (1993 ~ 2002 年)

年份	收益率 (%)	年份	收益率 (%)
1993	14.8	1998	9.2
1994	4.5	1999	3.8
1995	33.3	2000	13.1
1996	20.3	2001	1.6
1997	28.8	2002	-13.0

资料来源：AAIL。

利用表 3-28 中的信息，回答下列问题：

- (1) 计算 PRFDX 的偏度，结果保留到小数点后两位。
- (2) 基于第 1 题的解答，描述 PRFDX 收益率分布的形状。

⊖ 在式 (3-17) 中的 $n/[(n-1)(n-2)]$ 一项是对于小样本向下偏差的修正。
⊖ 要了解更多关于投资组合选择中偏度所发挥的作用，参见 Reilly 和 Brown (2003) 和 Elton 等 (2003) 以及这些书中所提及的其他相关资料。

解：

(1) 为了计算偏度，我们求解偏离均值偏差值的立方之和，然后将其除以标准差的立方，在将之前所得的结果乘以 $n/[(n-1)(n-2)]$ 。表 3-29 给出了所有的计算结果。

表 3-29 PRFDX 偏度的计算

年份	R_t	$R_t - \bar{R}$	$(R_t - \bar{R})^3$
1993	14.8%	3.16	31.554
1994	4.5%	-7.14	-363.994
1995	33.3%	21.66	10 161.910
1996	20.3%	8.66	649.462
1997	28.8%	17.16	5 053.030
1998	9.2%	-2.44	-14.527
1999	3.8%	-7.84	-481.890
2000	13.1%	1.46	3.112
2001	1.6%	-10.04	-1 012.048
2002	-13.0%	-24.64	-14 959.673
$n =$	10		
$\bar{R} =$	11.64%		
		总和 =	-933.064
$s =$	13.65%	$s^3 =$	2 543.302
		总和/ $s^3 =$	-0.366 9
		$n/[(n-1)(n-2)] =$	0.138 9
		偏度 =	-0.05

资料来源：AAIL

利用式 (3-17)，计算得：

$$S_k = \left[\frac{10}{(9)(8)} \right] \frac{-933.064}{13.65^3} = -0.05$$

在这个例子中，5 个偏差值为负，5 个偏差值为正。两个最大的正偏差，分别出现在 1995 年和 1997 年，差不多抵消了在 2002 年出现的很大的负偏差值和 2001 年出现的中等大的负偏差值（这两个熊市年份）。该结果表明偏度是一个非常小的负数。

(2) 基于这个小样本，基金年收益率的分布似乎近似是对称的（或者稍稍呈现出一点负偏）。偏离均值的负偏差和正偏差以相同的频数出现，大的正偏差差不多抵消了大的负偏差。

3.9 收益率分布的峰度

在之前几节中，我们讨论了如何根据偏度来确定一个收益率分布是否偏离正态分布。而另外一个确定收益率分布是否偏离正态分布的方法是通过观察是否有大部分的收益率集中于均值附近（呈现出尖峰）以及更多的收益率偏离均值的偏差极大（呈现出厚尾）。相对于正态分布，非正态分布具有很大比率的偏离均值偏差很小的收益率（更多的小意外）以及很大比率的偏离均值偏差很大的收益率（更多的大意外）。大部分投资者认为随着风险的增大，出现偏离收益率均值极端偏差的可能性将会增大。

峰度 (kurtosis) 是一个统计度量指标, 它告诉我们一个分布的峰值比正态分布更高还是更低。一个分布的峰值如果比正态分布更大, 我们就称其为尖峰的 (leptokurtic) (该词 lepto 是由希腊语“细长的” (slender) 而来); 一个分布的峰度如果比正态分布小, 那么其就被称为低峰的 (platykurtic) (该词 platy 是由希腊词汇宽阔的 (broad) 而来); 一个分布的峰度如果等于正态分布, 那么其就被称为中峰的 (mesokurtic) (该词 meso 是由希腊词汇中间的 (middle) 而来)。我们将出现更多极端的大意外的情形描述为尖峰中的一种。[Ⓐ]

图 3-8 描述了一个尖峰分布。相比正态分布, 它具有尖峰和厚尾的特征。

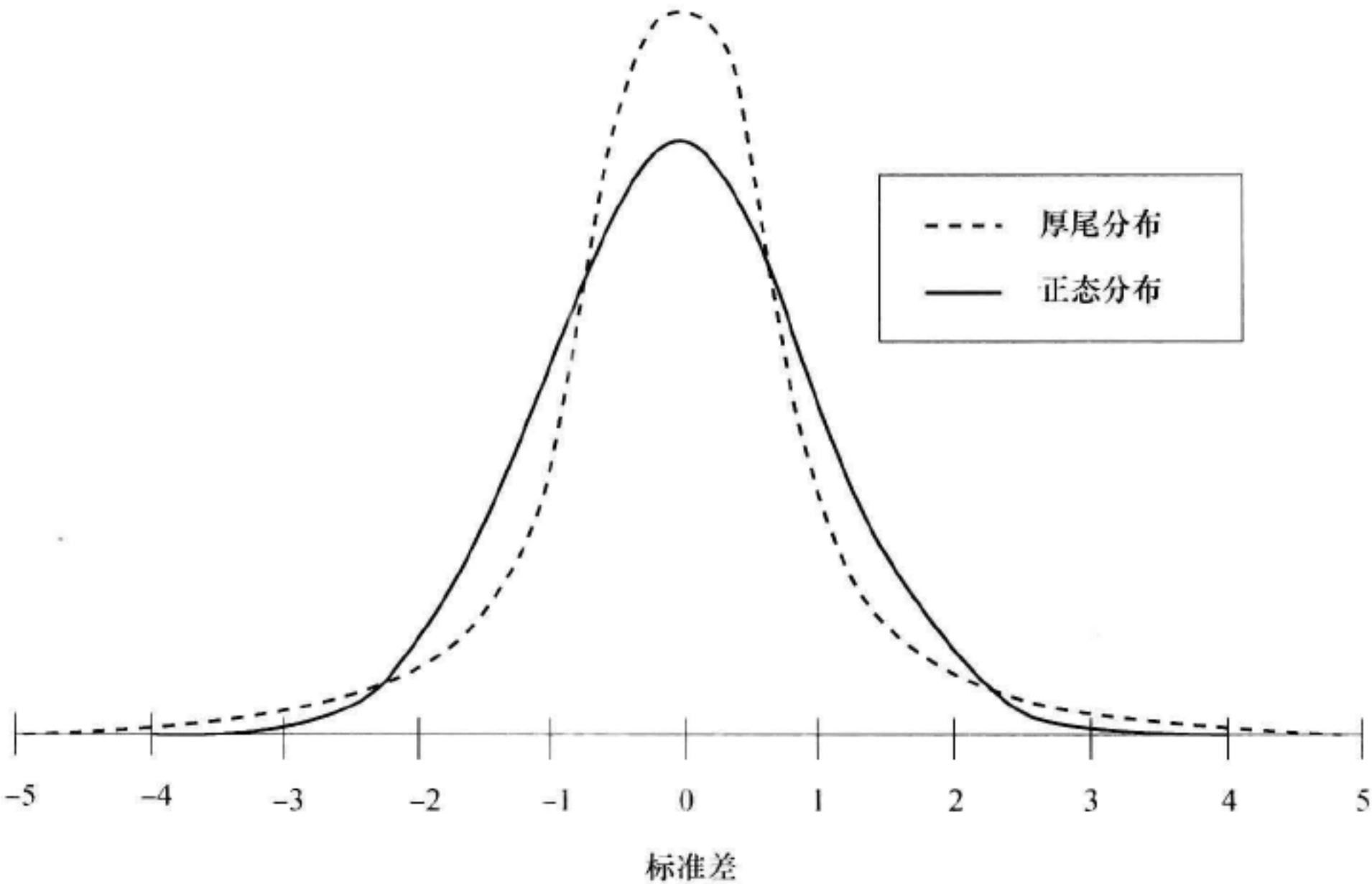


图 3-8 尖峰: 厚尾

资料来源: Reprinted with permission from *Fixed Income Readings for the Chartered Financial Analyst® Program*. Copyright 2000, Frank J. Fabozzi Associates, New Hope, PA.

峰度的计算涉及求解偏离均值偏差值四次方的平均值, 再将其除以标准差四次方的标准化计算。[Ⓑ]对于所有的正态分布来说, 峰度等于 3。许多统计软件会给出超额峰度 (excess kurtosis) (指峰度减去 3 之后的值) 的估计值。[Ⓒ]因此, 超额峰度描述了相对于正态分布峰度的值。一个正态分布或者其他中峰分布的超额峰度为 0。一个尖峰分布的超额峰度大于 0, 而一个低峰分布的超额峰度小于 0。一个具有正的超额峰度的收益率分布 (一个尖峰收益率分布) 相对于正态分布, 具有出现频率更高的极端偏离均值的偏差值。下面给出的是样本峰度的计算表达式。

- 样本超额峰度的计算公式 (sample excess kurtosis formula)。样本超额峰度 (sample excess kurtosis) 为:

Ⓐ 峰度被描述为由经常会出现极端情况造成的一种病态特征。
Ⓑ 该度量指标是没有量纲的。它总是为正, 因为偏差值被求了四次方。
Ⓒ 伊博森软件和其他一些软件包, 如微软 (Microsoft) 的 Excel, 将“超额峰度”简单地称为“峰度”。这里指出了这样一个事实, 即使用者必须熟悉其所使用的任何软件包所计算的描述性统计量的具体含义。

$$K_E = \left(\frac{n(n+1)}{(n-1)(n-2)(n-3)} \frac{\sum_{i=1}^n (X_i - \bar{X})^4}{s^4} \right) - \frac{3(n-1)^2}{(n-2)(n-3)} \tag{3-18}$$

式中， n 表示样本大小， s 表示样本标准差。

在式 (3-18) 中，样本峰度 (sample kurtosis) 是其第 1 项。注意到随着 n 的增大，式 (3-18)

近似等于 $\left(\frac{n^2}{n^3} \frac{\sum_{i=1}^n (X_i - \bar{X})^4}{s^4} \right) - \frac{3n^2}{n^2} = \frac{1}{n} \frac{\sum_{i=1}^n (X_i - \bar{X})^4}{s^4} - 3$ 。对于一个样本容量为 100 或者更大的从一个正态分布中所选取的样本来说，其样本超额峰度大于等于 1.0 可以被认为是异常大的。

我们发现大部分股票的收益率序列是尖峰的。如果一个收益率分布具有正的超额峰度 (尖峰) 且我们使用的统计模型无法解释其厚尾，那么我们将会低估其出现非常坏或者非常好结果的可能性。例如，1987 年 10 月 19 日标准普尔 500 的收益率偏离了其平均日收益率 20 个标准差。这一结果在正态分布中是可能出现的，但是其出现的概率几乎为 0。如果日收益率来源于一个正态分布，那么收益率偏离其均值 4 倍标准差及以上的出现概率为每 50 年；收益率偏离其均值 5 倍标准差及以上的出现概率为每 7 000 年。故 1987 年 10 月的收益率更可能来自一个比正态分布厚尾的分布。让我们看一下之前给出的表 3-27，其中，标准普尔 500 的月收益率序列具有很大的超额峰度，约为 9.5。这相对于正态分布来说就具有非常明显的厚尾。相反，年收益率序列只具有较小的负超额峰度 (大约 -0.2)。表中超额峰度的结果与人们的研究发现结果是一致的，即正态分布是对于美国年度持有期收益率的较好近似，而不是更短区间的收益率 (如月收益率)。[⊖]

下面的例子将举例说明我们之前已经考察过的两个共同基金之一的样本超额峰度的计算方法。

例 3-17 样本超额峰度的计算

例 3-16 中我们已经得到了普信股票收益基金的年收益率在 1993 ~ 2002 年区间近似服从对称分布的结论。那么该基金收益率分布的峰度是怎样的呢？表 3-28 (列示如下) 简要地重新给出了该基金的年收益率。

表 3-28 (续) 年收益率：普信股票收益基金 (1993 ~ 2002 年)

年份	收益率 (%)	年份	收益率 (%)
1993	14.8	1998	9.2
1994	4.5	1999	3.8
1995	33.3	2000	13.1
1996	20.3	2001	1.6
1997	28.8	2002	-13.0

资料来源：AAIL。

利用上面表 3-28 给出的这些信息，回答下列问题：

- (1) 计算 PRFDX 的样本超额峰度，结果保留到小数点后两位。
- (2) 基于对于第 1 题的回答，描述 PRFDX 收益率分布的形状到底是尖峰的、中峰的还是低峰的。

⊖ 更多细节参见 Campbell, Lo 和 MacKinlay (1997)。

解：

(1) 为了计算超额峰度，我们求解偏离均值偏差值四次方之和，然后将其除以标准差的四次方，在将之前所得的结果乘以 $n(n+1)/[(n-1)(n-2)(n-3)]$ 。这一计算得出了峰度。而超额峰度为峰度减去 $3(n-1)^2/[(n-2)(n-3)]$ 。表 3-30 给出了所有的计算结果。

表 3-30 PRFDX 偏度的计算

年份	R_t	$R_t - \bar{R}$	$(R_t - \bar{R})^4$
1993	14.8%	3.16	99.712
1994	4.5%	-7.14	2 598.920
1995	33.3%	21.66	220 106.977
1996	20.3%	8.66	5 624.340
1997	28.8%	17.16	86 709.990
1998	9.2%	-2.44	35.445
1999	3.8%	-7.84	3 778.020
2000	13.1%	1.46	4.544
2001	1.6%	-10.04	10 160.963
2002	-13.0%	-24.64	368 606.351
$n =$	10		
$\bar{R} =$	11.64%		
		总和 =	697 725.261
$s =$	13.65%	$s^4 =$	34 716.074
		总和/ $s^4 =$	20.098
		$n(n+1)/[(n-1)(n-2)(n-3)] =$	0.218 3
		峰度 =	4.39
		$3(n-1)^2/[(n-2)(n-3)] =$	4.34
		超额峰度 =	0.05

资料来源：AAIL

利用式 (3-18)，计算得：

$$\begin{aligned} K_E &= \left[\frac{110}{(9)(8)(7)} \frac{697\,725.261}{13.65^4} \right] - \frac{3(9)^2}{(8)(7)} \\ &= 4.39 - 4.34 \\ &= 0.05 \end{aligned}$$

(2) 基于样本超额峰度接近于 0 的结果，PRFDX 年收益率的分布表现为中峰的。由于偏度和超额峰度都接近于 0，PRFDX 的年收益率在该区间近似表现为正态分布。[⊖]

⊖ 如果我们可以根据样本容量 n 、样本偏度和样本超额偏度来计算一个 Jarque-Bera (JB) 正态性检验统计量，那么这将对非常有用。如果对于一个至少包含 30 个观测值的样本，其 $JB = n[(S_k^2/6] + (K_E^2/24)]$ 的值大于等于 6，我们就能在不超过 5% 犯错可能性的条件下，得出分布是否是正态分布的结论。在共同基金的例子中，我们只有 10 个观测值，而该检验只有在大样本的情况下才正确（根据要求，样本容量 n 要大于等于 30），因此，我们无法使用该检验。Gujarati (2003) 提供了更多关于该检验的相关内容。

3.10 使用几何平均和算术平均

在掌握了描述性统计量的概念之后，我们将阐述为什么几何平均适用于对于过去业绩的投资评估。我们也将探讨为何算术平均适用于对于未来预测的投资判断。

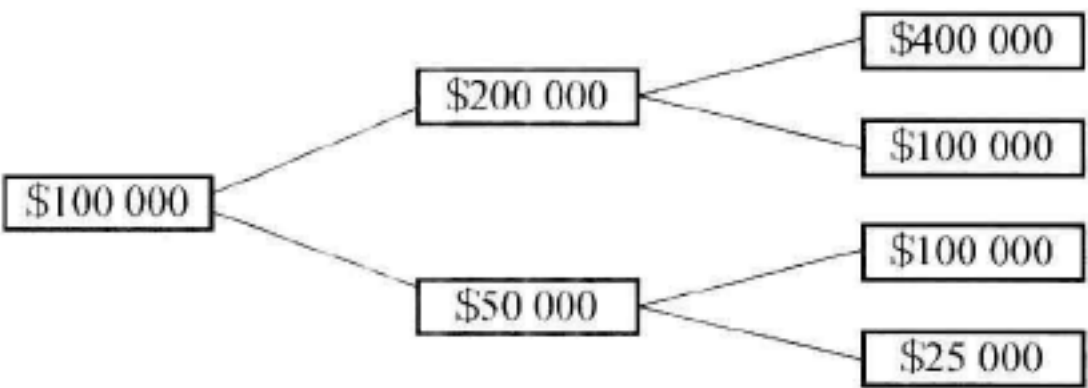
在报告历史收益率时，几何平均具有相当大的吸引力，因为它表示的是我们在每年已经获得的且与实际累计投资业绩相匹配的增长率或者收益率。例如，在我们简化的例 3-8 中，我们以 100 欧元购买了一份股票，在两年之后该股票价值 100 欧元，其中在第 1 年年末其价格为 200 欧元。几何平均数为 0%，它反映的显然是这两年间复合增长率。具体地来说，期末值是期初值乘以 $(1 + R_g)^2$ 。几何平均数是一个很不错的度量历史业绩的指标。

例 3-8 也表明了为何算术平均数会扭曲我们对于历史业绩的评价。在那个例子中，在这两年间的总收益显然应该为 0%。然而，由第 1 年 100% 的收益率和第 2 年 -50% 的收益率，我们所得到的算术平均数为 25%。正如我们之前所述的，算术平均数总是大于等于几何平均数。如果我们想要估计一个单期的平均收益率，我们应该使用算术平均数，因为算术平均数就是对于单期收益率的平均。然而如果我们想要估计多于一期的平均收益率，那么我们应该使用几何平均收益率，因为几何平均数能够描述多期总收益率间是如何相互联系的。

作为使用几何平均来进行业绩报告的一个推论，半对数（semilogarithmic）尺度要比算术尺度更适用于对历史业绩的作图。[⊖]在报告业绩的情形下，半对数尺度的图像是以算术尺度度量的时间作为横轴，而以对数尺度度量的投资额作为纵轴。纵轴上值与值之间的间隔是根据其对数值之间的间隔而得来的。假设我们想在纵轴上表示 1 英镑、10 英镑、100 英镑和 1 000 英镑这几个投资额。注意到每一个后续的值都要比前一个值大 10 倍，所以这四个值将会等间距地标注在纵轴之上，因为它们之间对数值的间隔近似为 2.30；即 $\ln 10 - \ln 1 = \ln 100 - \ln 10 = \ln 1000 - \ln 100 = 2.30$ 。在半对数的尺度下，纵轴上等间距的移动代表着数值等百分比的变化，故以固定复利增长率变化的数据在图上将被画成一条直线。一条向上弯曲的曲线将代表增长率随时间增长而变大的一组数据。图像在不同点的斜率可以相互比较，以判断它们相对增长率的情况。

另外，在报告历史业绩时，金融分析师们需要以一个前瞻性的视角来计算期望的股票风险溢价。为了达到这一目的，使用算术平均是合适的。

我们可以利用一项投资的未来现金流的例子来说明在前瞻性的视角下算术平均的应用。用几何平均和算术平均贴现未来现金流的本质差异在于对于不确定性的理解不同。假设一个拥有 100 000 美元的投资者面对以等概率出现的 100% 的收益率和 -50% 的收益率（这一情况可以表示为每期 100% 和 -50% 收益的 50/50 概率的树形图）。由第 1 期 100% 的收益率和第 2 期 -50% 的收益率，可以得到其几何平均收益率为 $\sqrt{2(0.5)} - 1 = 0$ 。



⊖ 更多信息详见 Campbell (1974)。

几何平均收益率 0% 给出了第二期期末财富的众数或者中位数，因此，在这个例子中，几何平均收益率能准确地预测出期末财富为 100 000 美元。但是，算术平均收益率能更好地预测出期末财富的算术平均值。在等概率出现 100% 和 -50% 收益率的条件下，考虑 4 个都可能出现的结果 400 000 美元、100 000 美元、100 000 美元和 25 000 美元（假设它们都实际发生了）。那么期末财富的算术平均值将为 $\$156\,250 = (\$400\,000 + \$100\,000 + \$100\,000 + \$25\,000)/4$ 。实际的收益率为 300%、0%、0% 和 -75%，于是两期算术平均收益率为 $(300 + 0 + 0 - 75)/4 = 56.25\%$ 。该算术平均收益率所预测的期末财富算术平均值为 $\$100\,000(1.5625) = \$156\,250$ 。注意到两期收益率 56.25% 说明每期的收益率为 25%，所以我们必须将最终期望财富 156 250 美元以 25% 的算术平均收益率进行贴现以反映现金流中的不确定性。

现金流或收益率中存在的确定性使得算术平均数比几何平均数要大。收益率越不确定，那么算术平均数和几何平均数之间存在的差别也就越大。几何平均收益率近似等于算术平均数减去其收益率方差的一半。[⊖]当收益率的方差为 0 或者不存在不确定性时，几何平均收益率和算术平均收益率将近似相等。但是现实世界中存在不确定性常使得算术平均收益率要大于几何平均收益率。例如，蒂姆森（Dimson）等人（2002）指出，1900 ~ 2000 年美国股票名义年收益率的算术平均收益率为 12%，标准差为 19.9%。而在他们的报告中，这些股票同期的几何平均收益率仅为 10.1%。由此我们可以看到，几何平均收益率近似等于算术平均收益率减去其收益率方差的一半： $R_g \approx 0.12 - \left(\frac{1}{2}\right)(0.199^2) = 0.10$ 。

⊖ 参见 Bodie, Kane 和 Marcus (2001)。



第

章

概率论中的一些概念

4.1 引言

所有的投资决策都是在有风险的环境下做出的。而在这样的环境中，能够帮助我们一致且逻辑地做出决策的工具都可以在概率论中找到。本章将介绍在建模过程中所需要使用的最基本的概率工具，并且对现实世界中许多与风险有关的问题加以回答。我们将举例说明如何应用这些工具来解决诸如预测投资经理的业绩表现、预测金融变量以及为债券进行合理定价（以使得它们能够很好地补偿债券持有者所承担的违约风险）等问题。我们所关注的都是实际的应用问题。我们所具体探讨的一些概念都是在投资研究和实践中最为重要的。其中的一个重要概念就是独立性，因为这个概念与收益率以及其他金融变量的可预测性密切相关。而另一个重要概念是期望值，因为分析师们在他们的分析和决策中始终考虑的是未来可能发生的情况。当然，分析师和投资者同样也要应对数据的波动性。作为投资学中重要的有关风险的概念，我们在本章中所介绍的是方差（或者是偏离期望值的偏差）。通过本章的学习，读者将学会使用投资组合的期望收益率和方差进行分析的具体技能。

概率论中的基本工具包括期望值和方差。它们将在本章的第2节中进行介绍。第3节将介绍协方差和相关系数（度量随机变量间相关程度的指标）以及计算投资组合期望值和方差的原理。在本章的最后，我们将讨论贝叶斯公式和对试验结果计数这两个问题。贝叶斯公式是在新信息的基础上，对于概率信念进行更新的一种计算方法。在不少领域中，包括最广泛使用的期权定价模型，都包含了对于试验结果的定义和计数的计算。本章的最后将讨论一下计数原理以及在概率计算中所使用的一些小诀窍。

4.2 概率、期望值和方差

在工作中，大部分分析师所需要用到的概率论概念和工具是相对较少和简单的，但是在应用它们时都需要我们仔细地去进行思考。这一节将通过股票和固定收益债券分析的例子来具体介绍与概率、期望值和方差应用相关的一些基本内容。

投资者非常关注收益率。风险资产的收益率是一个随机变量 (random variable) (一个量, 其出现的结果 (outcome) (可能值) 是不确定的)。例如, 一个投资组合可能会设定一个每年 10% 的收益率目标。这时, 投资经理所关注的就可能是明年出现收益率低于 10% 的可能性。这里的 10% 是“投资组合收益率”这个随机变量的一个特定值或者特定结果。虽然我们有时也会关心某一个结果, 但是通常我们会去关注一系列的结果: 而“事件”的概念就包含了前面我们所要关注的这两种情况 (单个结果和一系列结果)。

- **事件的定义。**一个事件 (event) 是指一组具体的实验结果。

我们可以将事件定义为单个试验结果。例如, 投资组合获得 10% 的收益率。我们将收益率低于 10% 的投资组合定义为一个事件来描述投资组合经理所关注的问题。第 2 个事件定义为所有大于等于 -100% (最差可能的收益率) 且小于 10% 的可能收益率。这个事件包括无限种结果。为了表述简单, 我们通常使用斜体的大写字母来表示所定义的事件。我们可以定义 A = 投资组合获得 10% 的收益率, 以及 B = 投资组合获得低于 10% 的收益率。

回到之前投资组合经理所关注的问题, 投资组合的收益率低于 10% 的可能性到底会有多少呢?

该问题的答案是一个概率 (probability): 一个介于 0 ~ 1 之间的数值, 其度量的是所关注事件会发生的可能性。如果投资组合收益率低于 10% 的概率是 0.4, 那么这就代表收益率低于 10% 事件有 40% 的可能性会发生。如果一个事件是不可能发生的, 那么它的概率就是 0。如果一个事件是必然发生的, 那么其概率就是 1。如果一个事件是不可能发生的或者是必然发生的, 那么它就不是随机的。所以, 0 ~ 1 之间的数包括了所有可能的概率取值。

概率具有两个性质, 这两个性质构成了概率的定义。

- **概率的定义。**概率的两个性质定义如下:
 1. 任何事件 E 的概率是一个介于 0 ~ 1 之间的实数: $0 \leq P(E) \leq 1$ 。
 2. 任何一组两两互斥且穷举的事件序列的概率之和为 1。

我们用 P 和之后的括号来表示“(括号内的事件)发生的概率”, 如 $P(E)$ 表示“事件 E 发生的概率”。我们也能将 P 视为一个满足性质 1 和性质 2 的给予事件分配数值的函数或规则。

在之前的定义中, 两两互斥 (mutually exclusive) 这个术语指的是两个事件不可能同时发生; 穷举 (exhaustive) 意味着 (事件序列中的) 事件包含了所有可能发生的结果。事件 A = 投资组合获得 10% 的收益率和事件 B = 投资组合获得低于 10% 的收益率是互斥的, 因为事件 A 和 B 不能同时发生。例如, 出现 8.1% 的收益率表明 B 发生, 而 A 没有发生。虽然事件 A 和事件 B 是互斥的, 但是它们不是穷举的, 因为它们没有包括所有可能的结果 (如 11% 的收益率)。假设我们定义第 3 个事件 C = 投资组合获得高于 10% 的收益率。那么显然, 事件 A 、 B 和 C 就是两两互斥且是穷举的事件序列。于是 $P(A)$ 、 $P(B)$ 和 $P(C)$ 都是 0 ~ 1 中的一个数字, 且 $P(A) + P(B) + P(C) = 1$ 。

两两互斥且穷举的事件序列中最基本的一个事件就是由随机变量所有可能的不同结果组成的集合。如果我们知道了这个集合以及所有结果会发生的概率 (随机变量的概率分布), 那么我们对随机变量就有了一个完整的描述, 并且我们可以给出任一个我们所要描述事件的概率值。[⊖]任何事件的概率等于该给定的事件中每个可能出现结果的概率之和。假设我们所关注的事件为 D =

⊖ 在常用概率分布一章中, 我们将介绍在投资应用中最为常用的一些概率分布。

投资组合获得高于无风险收益率的收益率，并且我们已知投资组合收益率的概率分布。假设无风险利率是4%。为了计算 $P(D)$ (事件 D 的概率)，我们将对于满足该事件定义的所有可能结果发生的概率进行加总。即我们将加总所有投资组合收益率超过4%事件的概率值。

之前，为了介绍概念，我们假设了投资组合收益率低于10%的概率为0.40。除此之外，没有做其他假设。我们也论述了利用试验结果概率分布来计算事件发生概率的方法。但是，我们并没有说明如何对于概率分布进行估计。利用不准确的概率进行实际的金融决策可能会造成严重的后果。那么，在实际中，我们应如何来估计概率呢？该问题本身就是一个值得研究的广阔领域，但是概括起来，估计概率主要有3种方法。在投资过程中，我们经常把由历史数据所获得的相对发生频数作为事件发生概率的估计值。该方法产生的是一个经验概率 (empirical probability)。例如，阿弥哈德 (Amihud) 和李 (Li) (2002) 指出在他们所使用的16 189个1962~2000年间纽约证券交易所和美国证券交易所中交易股票的红利变化值的样本中，14 911个是增加的，而1 278个是减少的。因此，美国股票减少的红利变化值的经验概率约为 $1\,278/16\,189 = 0.08$ 。在本章中，我们将在每次经验概率出现的地方指出它们。

在时间上表现出稳定关系的经验概率是准确的。我们不能计算那些没有历史数据记录的事件的经验概率，也不能计算出一个非常罕见的突发事件的准确经验概率。因此，我们可能在一些情况中对经验概率进行调整，使其能够解释我们所能感知的变动关系。也可能在另一些情况中，我们根本没有经验概率可以利用。我们还可能会不基于任何特定的数据，给出我们个人对于概率的主观判断。我们在这三种情况下所获得的概率都是主观概率 (subjective probability)，一种基于个人或者主观判断而获得的概率值。主观概率在投资过程发挥着非常重要的作用。投资者在确定资产价格的买卖决策过程中，经常关注主观概率。主观概率将会在本章多次出现，特别是在我们讨论贝叶斯公式的时候。

在较小范围的一些定义良好的问题中，我们有时可以通过对于问题的逻辑判断推导出概率。所得到的概率是先验概率 (a priori probability)，它是指基于逻辑分析而非观测所形成的概率或个人主观判断的概率。我们将在例4-6中使用这种概率。我们之后讨论的计数方法在计算先验概率过程中非常重要。因为先验概率和经验概率一般不会因人而异，它们一般被归为客观概率 (objective probabilities) 一类。

在商业或者其他一些领域，我们经常会遇到以优比来表述的概率。例如，“ E 发生的优比”或者“ E 不发生的优比”。例如，在2003年中期，分析师们对于多伦多证券交易所的一家上市公司在2004财年EPS的预测值在3.98~4.25加元之间。但是，一位分析师声称该公司的EPS超过其最高估计值4.25加元的优比为1:7。而另一位分析师认为该事件不发生的优比应该为15:1。那么，这些论断所表示的该公司EPS超过估计最高值的概率是多少？我们下面对以优比表示的概率进行解释。

- 以优比 (odds) 表示的概率。给定一个概率 $P(E)$,

1. E 发生的优比为 $= P(E)/[1 - P(E)]$ 。 E 的优比是 E 发生的概率除以1减去 E 发生的概率。给定 E 发生的优比为“ $a:b$ ”，那么这意味着 E 发生的概率为 $a/(a+b)$ 。

在例子中，对于2004财年该公司EPS超过4.25加元事件发生的优比为1:7，所以这表示该分析师认为该事件发生的概率为 $1/(1+7) = 1/8 = 0.125$ 。

2. E 不发生的优比 $= [1 - P(E)]/P(E)$ ，这是 E 发生优比的倒数。给定 E 不发生的优比为

“ $a:b$ ”，那么这意味着 E 发生的概率为 $b/(a+b)$ 。

故对于该公司 2004 财年 EPS 超过 4.25 加元事件不发生的优比为 15:1 的表述与该事件发生的概率为 $1/(1+15) = 1/16 = 0.0625$ 的信念是一致的。

为了进一步解释一个事件发生的优比，假设 $P(E) = 1/8$ ，那么 E 发生的优比为 $(1/8)/(7/8) = (1/8)(8/7) = 1/7$ ，或者“1:7”。对于每次事件 E 的发生，我们预计会有 7 次不发生的情况出现；因此，在 8 次的试验中，我们预计事件 E 会发生一次，故其概率为 $1/8$ 。而在赌博中，我们通常以某事件不发生的优比（如同表述 2）来进行表达。对于 E 不发生的优比“15:1”（隐含着 E 发生的概率为 $1/16$ ），对于 E ，如果成功的话，1 美元的赌资的收益为 15 美元的奖金和 1 美元的下注本金。我们能计算下注的期望收益如下：

赢: 概率 = $1/16$; 收益 = 15 美元
输: 概率 = $15/16$; 收益 = -1 美元
期望收益 = $(1/16)(\$15) + (15/16)(- \$1) = \$0$

如果优比（概率）是准确的话，那么对于两种赌博结果的每一种以其发生的概率进行加权，我们就得到该赌博的期望收益为 0 美元。

例 4-1 由不一致的概率所得到的收益

你现在正在对同一行业的两家公司的普通股进行考察，而该行业将在下周有一个重要的反垄断决定会公布。第 1 家公司，SmithCo 公司将会从政府的决定中获益，因为它所涉及的兼并活动并没有面临反垄断的障碍。你认为 SmithCo 公司的股价对于该决定做出利好反应的概率为 0.85。而第 2 家公司，Selbert 公司也将会从预先的管理中获得同等大小的收益。但是出人意料地，你认为 Selbert 股票仅会以 0.50 的可能性对于该利好消息做出反应。假设你的分析都是正确的，那么怎样的投资策略能够从该定价差异中获益？

让我们考虑一下逻辑上的可能性。一种可能是 Selbert 公司股价以 0.50 概率做出反应的判断是正确的。在这种情况下，Selbert 公司的估值是正确的，而 SmithCo 公司被高估了，因为它的当前股价高估了对于预先管理利好反应的概率。第 2 种可能是 0.85 的概率是正确的。在这种情况下，SmithCo 公司股票的估值是正确的，而 Selbert 股票建立在对于利好的决定以较低概率反应基础上的估值被低估了。你可以在表 4-1 中将这两种情况以图表形式表现出来。

0.50 概率这一栏表示 Selbert 公司股票比 SmithCo 公司股票估值要更准确。而同样地在 0.85 的概率下，Selbert 股票的估值也更准确。因此，SmithCo 公司股票相对于 Selbert 公司股票是被高估了。

你的投资行为取决于你对于你的分析的自信程度以及你所面临的投资约束（例如对于卖空股票的限制）。[⊖] 一个保守的策略为买 Selbert 公司股票而降低或者清空目前 SmithCo 公司股票的仓位。最为激进的策略为卖空 SmithCo 公司股票（相对高估），而与此同时买入 Selbert 公司股票（相对低估）。这一策略被俗称为配对套利交易（pairs arbitrage trade），指的是对于两个表现密切相关股票的一种交易，其涉及卖空其中一只股票同时买入另一只股票。

表 4-1 投资问题的工作表

	“预先”决定的真实概率	
	0.50	0.85
SmithCo 公司	股价高估	股价估值正确
Selbert 公司	股价估值正确	股价低估

⊖ 卖空或者卖空股票指卖出借来的股票，以期在未来以更低的价格再买入它们。

SmithCo 公司和 Selbert 公司的股票价格（对于政府决定）反应的概率是不一致的。根据一个投资中最重要的概率结果荷兰赌定理（Dutch book theorem）[Ⓐ]，指的是不一致的概率会产生获利机会。在我们的例子中，投资者会通过他们利用不一致概率的买卖决策，来消除获利机会和概率的不一致性。

为了理解投资环境中概率的含义，我们需要区分两种概率：无条件概率和条件概率。无条件概率和条件概率都满足我们之前对于概率的定义。但是它们的计算和估计是不同的，而且它们有各自不同的含义，它们会针对不同的问题给出答案。

对于直接的问题“这个事件 A 的概率是多少？”所给出的概率回答是无条件概率（unconditional probability），记为 $P(A)$ 。无条件概率也常被称为边际概率（marginal probabilities）。[Ⓑ]

假如问题是“股票获得的收益率高于无风险收益率（事件 A ）的概率是多少？”，那么该问题的回答就是无条件概率，它可以看做两个量的比值。分子是股票收益率高于无风险收益率的概率之和。假设该和为 0.70，分母是 1，它是所有可能收益率的概率。故该问题的结果为 $P(A) = 0.70$ 。

将问题“事件 A 发生的概率”与问题“给定事件 B 发生的情况下，事件 A 发生的概率”进行对比，由后一个问题给出的概率解答是条件概率（conditional probability），记为 $P(A|B)$ （读作：“给定事件 B ，事件 A 的概率”）。

假设我们想了解给定股票获得正收益（事件 B ）条件下，股票获得超过无风险收益率的收益（事件 A ）的概率。在“给定”这个词之后，我们将结果限制为大于 0% 的那些收益率——相对于之前的无风险概率问题，这是一个新的事件。条件概率同样可以计算为两个量之比。分子是股票收益率高于无风险收益率的概率之和；在这个例子中，分子与之前无条件概率例子中的一样，为我们给定的 0.70。然而，分母由 1 变为了所有高于 0% 收益率的概率之和。假设这个数值是 0.80，一个比 0.70 大的数值，因为在 0 和无风险收益率之间的收益率有正的发生概率，于是 $P(A|B) = 0.70/0.80 = 0.875$ 。如果我们观测到股票获得了一个正收益，那么这个收益率高于无风险收益率的概率要高于无条件概率（该事件的概率没有给定任何信息）。这个结果是直观的。[Ⓒ]让我们回顾一下，一个无条件概率是一个没有任何限制事件的概率，它甚至可能被认为是一个独立的概率。而相反，一个条件概率是一个在给定另一个事件发生情况下的事件发生概率。

在讨论计算概率的方法时，我们曾经给出一个产生股利下降变化概率的实际估计。那个概率是一个无条件概率。给定其他的关于公司特征的信息，那么投资者能否调整那个估计值？投资者一直在寻找新的信息以帮助其改善预测。从数学表达式的角度来看，他们试图利用基于相关的信息或者事件的概率来形成对于未来的看法。投资者不会忽略有用的信息，他们会通过调整他们的概率来反映这些信息。因此，条件概率的概念（我们将在下面更加详细地分析），与之后将会深入讨论的其他相关概念都是在投资分析以及金融市场中非常重要的。

为了准确地给出条件概率的定义，我们首先需要引入联合概率的概念。假如我们问这样的问题：

Ⓐ 该定理的名称来源于赌博中的术语。假设一个人基于 X 发生相对于 X 不发生的优比为 10:1，下注 100 美元。而之后，他又基于 X 不发生相对于 X 发生的优比为 1:1，故又下注 600 美元。无论 X 的结果如何，这个人将获得无风险的收益（如果 X 发生的话，收益等于 400 美元；而如果 X 不发生的话，收益等于 500 美元），因为隐含概率是不一致的。该人被称为在 X 上下了一个荷兰赌注。Ramsey (1931) 提出了不一致概率的问题。该问题同样可以参见 Lo (1999)。

Ⓑ 在分析表中表述的概率时，无条件概率通常出现在表格的底部或者边上，因此就有了术语“边际概率”。因为可能与经济学中所使用的边际（大概表示为一种增量）一词相混淆，我们在讨论中使用无条件概率这一说法。

Ⓒ 在这个例子中，条件概率比无条件概率要大。然而，一个事件的条件概率可能大于、等于或者小于无条件概率，这取决于具体的情况，例如，给定股票获得负收益条件下，股票获得超过无风险收益率的收益的概率为 0。

“事件 A 和 B 共同发生的概率是多少?” 那么该问题的答案就是一个联合概率 (joint probability), 记为 $P(AB)$ (读做事件 A 和 B 共同发生的概率)。如果我们认为事件 A 的概率和事件 B 的概率是由一个或超过一个随机变量的可能结果组成的集合, 那么事件 A 和 B 的联合概率就是这两个事件中共同含有的那些结果发生的概率之和。例如, 考虑两个事件: 股票获得超过无风险收益率的收益 (事件 A) 以及股票获得正的收益 (事件 B)。事件 A 的结果包含于事件 B 的结果之中 (或者说事件 A 的结果是事件 B 结果的一个子集), 所以 $P(AB)$ 等于 $P(A)$ 。我们现在可以用一个计算公式来给出条件概率的一个正式的定义。

- **条件概率的定义。** 给定事件 B 已经发生时, 事件 A 发生的条件概率等于事件 A 和 B 共同发生的联合概率除以事件 B 发生的概率 (假设其不等于 0)。

$$P(A | B) = P(AB) / P(B), P(B) \neq 0 \tag{4-1}$$

有时, 我们已知条件概率 $P(A | B)$, 而想要去了解联合概率 $P(AB)$, 我们可以从概率的乘法法则 (multiplication rule for probabilities) 获得联合概率。该法则是将式 (4-1) 重新整理而得到的。

- **概率的乘法法则。** 事件 A 和 B 的联合概率可以表述成:

$$P(AB) = P(A | B)P(B) \tag{4-2}$$

式 (4-2) 说明事件 A 和 B 共同发生的联合概率等于给定事件 B 发生条件下事件 A 发生的概率乘以事件 B 发生的概率。因为 $P(AB) = P(BA)$, 所以 $P(AB) = P(BA) = P(B | A)P(A)$ 和式 (4-2) 是等价的。

例 4-2 条件概率和共同基金业绩的可预测性 (1)

卡恩 (Kahn) 和路德 (Rudd, 1995) 利用一个包含 300 个对美国国内股票进行积极管理的共同基金的样本, 对于历史业绩是否可以预测未来业绩的问题进行了检验。其中他们所用的一个方法涉及计算每个基金相对于一组风格指数的超额表现 (术语“风格”指的是成长型/价值型以及大市值/中市值/小市值这些股票特征)。为每个基金建立一个风格基准 (一个与基金风格相匹配的对照投资组合) 之后, 卡恩和路德计算了两期的基金选股的收益率。他们定义的选股收益率为基金的收益率减去基金所对应的风格基准的收益率。第一个期间为 1990 年 10 月至 1992 年 3 月。对选股收益率进行排序, 前 50% 的基金标记为赢家; 最后 50% 的基金标记为输家。基于下一期 (1992 年 4 月~1993 年 9 月) 的选股收益率, 同样将前 50% 的基金标记为赢家, 而最后 50% 的基金标记为输家。表 4-2 中摘录了他们的部分结果。例如, 赢家—赢家条目表示 150 家中的 79 家在第 1 期标记为赢家的基金在第 2 期仍标记为赢家 ($52.7\% = 79/150$)。注意表中 4 个条目中括号内的数字可以视为条件概率。

表 4-2 股票选股收益率

期间 1: 1990 年 10 月至 1992 年 3 月
期间 2: 1992 年 4 月至 1993 年 9 月

	期间 2 的赢家	期间 2 的输家
期间 1 的赢家	79 (52.7%)	71 (47.3%)
期间 1 的输家	71 (47.3%)	79 (52.7%)

注: 条目中的是基金数目 (括号中的是占行中基金总数的百分比)。
资料来源: 卡恩和路德 (1995) 中的表 3。

基于表 4-2 中的数据, 回答下列问题:
(1) 说明所需的用来定义 4 个条件概率的 4 个事件。

(2) 用 $P(\text{这个事件} | \text{那个事件}) = \text{数值}$ 的形式来表述表中 4 个条目的条件概率。

(3) 第 2 问的条件概率是经验概率、先验概率还是主观概率？

(4) 利用表中的信息，计算一个基金在第 1 期和第 2 期均为输家事件的概率。（注意：因为在每期有 50% 的基金被归为输家，在每期基金被标记为输家的无条件概率为 0.5。）

解：

(1) 用来定义条件概率的 4 个所需的事件为：

基金在第 1 期为赢家

基金在第 1 期为输家

基金在第 2 期为赢家

基金在第 2 期为输家

(2) 对于第 1 行，

$P(\text{基金在第 2 期为赢家} | \text{基金在第 1 期为赢家}) = 0.527$

$P(\text{基金在第 2 期为输家} | \text{基金在第 1 期为赢家}) = 0.473$

对于第 2 行，

$P(\text{基金在第 2 期为赢家} | \text{基金在第 1 期为输家}) = 0.473$

$P(\text{基金在第 2 期为输家} | \text{基金在第 1 期为输家}) = 0.527$

(3) 这些概率是由数据计算而来的，所以他们是经验概率。

(4) 概率的估计值为 0.264。事件 A 为一个基金在第 2 期为输家，事件 B 为基金在第 1 期为输家，而事件 AB 为基金在第 1 期和第 2 期均为输家。由表 4-2，有 $P(A | B) = 0.527$ 以及 $P(B) = 0.50$ 。因此，利用式 (4-2)，我们计算得：

$$P(AB) = P(A | B)P(B) = 0.527 \times (0.50) = 0.2635$$

或者概率约为 0.264。

当我们有两个感兴趣的事件 A 和 B ，我们经常想知道事件 A 或者事件 B 发生的概率。这里的“或者”是指并集，并集的意思为或者事件 A 发生或者事件 B 发生或者事件 A 和 B 共同发生。换句话说，事件 A 或者事件 B 发生的概率是这两个事件至少有一个发生的概率。该概率的计算要用到概率的加法法则 (addition rule for probabilities)。

- 概率的加法法则。给定事件 A 和 B ，事件 A 和 B 中至少有一个发生的概率等于事件 A 发生的概率加上事件 B 发生的概率减去事件 A 和 B 同时发生的概率。

$$P(A \text{ 或者 } B) = P(A) + P(B) - P(AB) \quad (4-3)$$

如果我们认为事件 A 的概率和事件 B 的概率是由一个或超过一个随机变量的可能结果组成的集合，那么计算事件 A 或者事件 B 发生的概率的第 1 步是加总事件 A 中所有结果发生的概率以得到 $P(A)$ 。如果事件 A 和 B 具有某些共同的结果，那么我们将 $P(A)$ 加上 $P(B)$ ，就会使得这个共同结果的概率计算两次。所以我们加入 $P(A)$ 中的量应为 $P(B) - P(AB)$ ，这一数量为事件 B 中除去在计算 $P(A)$ 时已经计算的那些共同结果的余下结果的概率。图 4-1 描述了这一计算过程：我们通过减去 $P(AB)$ 避免了对于 A 和 B 交集部分的重复计算。作为一个计算的例子，如果 $P(A) = 0.50$ ， $P(B) = 0.40$ 以及 $P(AB) = 0.20$ ，那么 $P(A \text{ 或者 } B) = 0.50 + 0.40 - 0.20 = 0.70$ 。只有在两个事件 A 和 B 互斥时，即 $P(AB) = 0$ 时， $P(A \text{ 或者 } B) = P(A) + P(B)$ 的表述才是正确的。

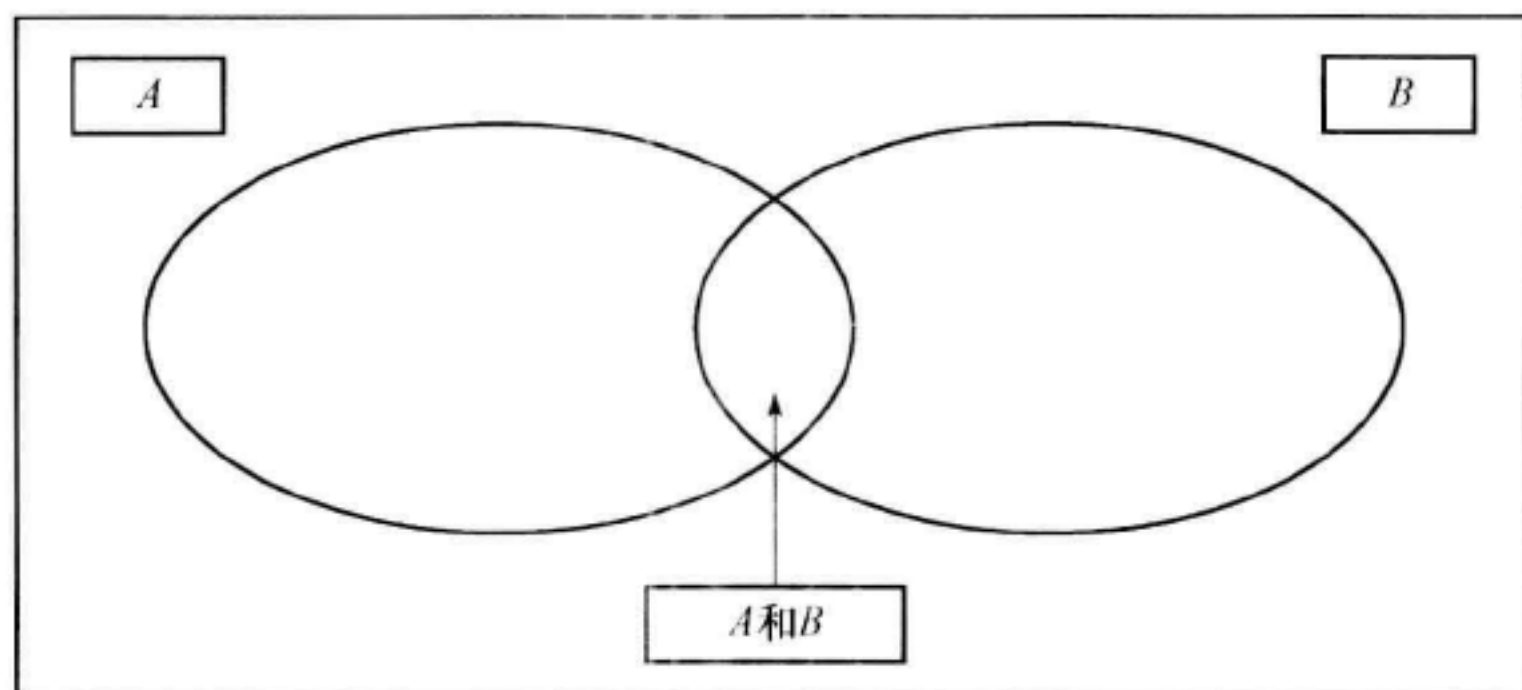


图4-1 概率的加法法则

下面的例子说明了利用目前的这些概率法则将会给我们带来多少有用的信息。

例4-3 执行限价指令的概率

你有两个已经发出的购买同一只股票的限价指令。一个报价购买股票的限价指令是指一个在该报价价格或者其价格之下购买股票的指令。许多卖方，包括你所使用的因特网服务，会提供给定当前股价和限定价格条件下，限价指令在一个给定区间内执行的估计概率。一个购买的指令（指令1）设定的限定价格为10美元，该指令在1小时内会被执行的概率为0.35。第2个购买指令（指令2）设定的限定价格为9.75美元；其在相同的1小时内会被执行的概率为0.25。

（1）指令1或者指令2将被执行的概率为多少？

（2）在给定指令1被执行的条件下，指令2将被执行的概率为多少？

解：

（1）概率为0.35。两个已知概率为 $P(\text{指令1被执行}) = 0.35$ 以及 $P(\text{指令2被执行}) = 0.25$ 。注意到如果指令2被执行，那么指令1一定也会被执行，因为价格一定会突破10美元而到达9.45美元。因此，

$$P(\text{指令1被执行} | \text{指令2被执行}) = 1$$

和

$$P(\text{指令1被执行且指令2也被执行}) = P(\text{指令1被执行} | \text{指令2被执行})$$

$$P(\text{指令2被执行}) = 1(0.25) = 0.25$$

为了回答该问题，我们会使用概率的加法法则：

$$\begin{aligned} P(\text{指令1被执行或者指令2被执行}) &= P(\text{指令1被执行}) + P(\text{指令2被执行}) \\ &\quad - P(\text{指令1被执行且指令2也被执行}) \\ &= 0.35 + 0.25 - 0.25 = 0.35 \end{aligned}$$

注意到指令2被执行的结果是指令1被执行结果的一个子集。在你计算出指令1被执行的概率之后，你也就算出了指令2也被执行的概率。因此，该问题的答案是指令1被执行的概率为0.35。

（2）如果第1个指令被执行，那么指令2被执行的概率为0.714。在第1题的解答中，你求得了 $P(\text{指令1被执行且指令2也被执行}) = P(\text{指令1被执行} | \text{指令2被执行})P(\text{指令2被执行}) = 1(0.25) = 0.25$ 。一个表述联合概率的等价方法在这里非常有用：

$P(\text{指令 1 被执行且指令 2 也被执行}) = 0.25 = P(\text{指令 2 被执行} | \text{指令 1 被执行})P(\text{指令 1 被执行})$

因为 $P(\text{指令 1 被执行}) = 0.35$ 是已知的, 所以你只有一个等式是未知的:

$$0.25 = P(\text{指令 2 被执行} | \text{指令 1 被执行})(0.35)$$

你可以得到 $P(\text{指令 2 被执行} | \text{指令 1 被执行}) = 0.25/0.35 = 5/7$, 或者约为 0.714。你也可以利用式 (4-1) 得到这个答案。

投资分析师非常感兴趣的是独立和不独立的概念。这些概念涉及一些基本的投资问题, 诸如那些金融变量对于投资分析是否有用、资产回报率是否可预测以及是否可以基于历史记录选出最好的投资经理。

当一个事件的发生不影响另一个事件发生的概率, 那么这两个事件就是独立的。

- **独立事件的定义。**两个事件 A 和 B 是独立的 (independent) 当且仅当 $P(A | B) = P(A)$, 或者等价地, $P(B | A) = P(B)$ 。

当两个事件不是独立的时候, 那么它们就是不独立的 (dependent): 一个事件发生的概率与另一个事件的发生有关。如果我们试图预测一个事件, 关于该事件相关事件的信息可能是有用的, 而与该事件独立的事件的信息是没有用的。

当两个事件是独立的时候, 概率的乘法法则, 式 (4-2) 将得以简化, 因为那个等式中的 $P(A | B)$ 将等于 $P(A)$ 。

- **独立事件的乘法法则** (multiplication rule for independent events)。当两个事件独立时, 事件 A 和 B 的共同概率等于事件 A 和 B 各个概率的乘积。

$$P(AB) = P(A)P(B) \quad (4-4)$$

因此, 如果我们对于两个概率分别为 0.75 和 0.50 的独立事件感兴趣时, 这两个事件同时发生的概率为 $0.375 = 0.75 (0.50)$ 。独立事件的乘法法则可以推广到多于两个事件的情形; 例如, 如果事件 A 、 B 和 C 是相互独立的事件, 那么 $P(ABC) = P(A)P(B)P(C)$ 。

例 4-4 BankCorp 的每股收益 (1)

作为一名银行业分析师工作的一部分, 你需要建立模型对于你所研究银行的每股收益进行预测。今天, 你研究的是 BankCorp。历史记录表明 55% 季度 BankCorp 的 EPS 是连续增长的, 而 45% 季度的 EPS 是连续降低或者是持续不变的。[⊖]你的分析就在这一点上, 假设 EPS 前后的变化间是相互独立的。

2004 年第二季度的每股收益 (即 2004 年第二季度的 EPS) 比 2004 年第三季度的 EPS 要大。

(1) 2004 年第三季度的 EPS 将大于 2004 年二季度 EPS 的概率是多少 (EPS 的一个正的连续变化)?

(2) 在之后两个季度 EPS 降低或者保持不变的概率是多少?

⊖ 对于季度 EPS 顺序上的比较是指与前一个季度相比。一个顺序上的比较和与前一年同季度间的比较 (另外一种常见的比较类型) 是不同的。

解:

- (1) 在独立的假设下, 2004 年第三季度 EPS 将会大于 2004 年第二季度 EPS 的概率为正变化的无条件概率 0.55。2004 年二季度 EPS 大于 2004 年一季度 EPS 的信息是没用的, 因为下个季度的 EPS 变化与前一个季度的变化相互独立。
- (2) 该概率为 $0.2025 = 0.45(0.45)$ 。

接下来的例子说明了即使每个条件本身并不是非常的严格, 但是要满足独立的所有条件却是非常困难的。

例 4-5 为投资筛选股票

你已经建立了一个股票筛选的要求——一组选择股票的条件。你的投资集 (你可以从中进行选择的一组股票) 为罗素 1000 指数, 该指数由 1000 个美国大市值股票组成。你筛选的要求包括了该选择问题的多个方面; 你认为每个要求间是相互独立的, 或者说是近似相互独立的。

筛选要求	满足要求的罗素 1000 股票的占比	筛选要求	满足要求的罗素 1000 股票的占比
第一估值要求	0.50	公司盈利能力要求	0.55
第二估值要求	0.50	公司融资能力要求	0.67
分析师荐股要求	0.25		

你预计会有多少股票通过你的筛选?

只有 1 000 只股票中的 23 只会通过你的筛选。如果你定义 5 个事件——股票通过第一估值要求、股票通过第二股指要求、股票通过分析师的荐股要求、股票通过公司盈利能力要求、股票通过公司融资能力要求 (不妨分别表示为事件 A、B、C、D 和 E)。那么在独立的条件下, 一只股票通过所有 5 个筛选要求的概率为:

$$P(ABCDE) = P(A)P(B)P(C)P(D)P(E) = (0.50)(0.50)(0.25)(0.55)(0.67) = 0.023\ 031$$

虽然 5 个要求中只有一个是相对非常严格的 (最严格的是 25% 的股票通过), 但是一只股票通过所有 5 个筛选要求的概率只有 0.023 031, 或者约 2%。候选投资列表中的规模为 $0.023\ 031(1\ 000) = 23.031$, 或者 23 只股票。

投资经理和其客户非常关注的一个问题是历史业绩的表现记录对于识别重复出现的赢家与输家是否有效。下面的例子说明了该问题是与独立的相关概念有关的。

例 4-6 条件概率和共同基金业绩的可预测性 (2)

如例 4-2 中所介绍的, 卡恩 (Kahn) 和路德 (Rudd, 1995) 的研究目的为探讨共同基金赢家和输家是否会重复出现的问题。如果一个基金在某一期为赢家或者输家的状态与其在下一期是否为赢家的情况间是独立的, 那么对于历史业绩排名的实际价值就是值得质疑的。利用例 4-2 表格中所定义的 4 个事件, 我们能定义如下事件来描述共同基金业绩可预测性的问题:

- 基金在第 1 期是赢家而在第 2 期也是赢家
- 基金在第 1 期是赢家而在第 2 期是输家

基金在第1期是输家而在第2期是赢家

基金在第1期是输家而在第2期也是输家

在例4-2的第4题中,你已经计算出

$$P(\text{基金在第2期是输家且在第1期也是输家}) = 0.264$$

如果在第1期的排名与第2期的排名间是独立的,那么你将预计 $P(\text{基金在第2期是输家而在第1期也是输家})$ 是多少?回答为经验概率0.264。

由独立事件的概率乘法法则, $P(\text{基金在第2期是输家而在第1期也是输家}) = P(\text{基金在第2期是输家})P(\text{基金在第1期是输家})$ 。因为在每期有50%的基金被归为输家,所以一个基金在某一期被标为输家的无条件概率为0.50。因此, $P(\text{基金在第2期是输家})P(\text{基金在第1期是输家}) = 0.50(0.50) = 0.25$ 。如果一个基金在第1期为输家的状态与之后一期是否为输家的状态间是独立的,那么我们可以得到 $P(\text{基金在第2期是输家而在第1期也是输家}) = 0.25$ 。这个概率是一个先验概率,因为它通过对于该问题的分析推导出来的。你同样可以将上面所描述的4个事件分别定义为不同的类别,如果基金被随机分配到这4个类别之中,那么基金在第1期是输家而在第2期也是输家的概率应该为1/4。如果在第1期和第2期的分配不是相互独立的,那么该基金的分类将不会是随机的。经验概率0.264只是稍稍高于0.25。这非常微小的可预测性是否来源于偶然情况?卡恩和路德检验的结果表明,如果第1期和第2期间的排名相互独立的话,有35.6%的可能性会观测到表中的数据。

在投资中,有关一个事件(或特征)是否会提供另一个事件(或特征)信息的问题存在于时间序列情形中(不同时点间),也存在于横截面数据情形中(给定一个时间点上不同个体间)。例4-4和例4-6说明了时间序列情形中的独立性。例4-5说明了横截面数据情形中的独立性。独立/不独立关系也经常在这两种情形中利用回归分析进行研究。这一方法我们将在之后的章节中进行介绍。

在许多实际问题中,我们会这样逻辑地分析一个问题:我们会明确列出认为会影响我们所关注事件可能性的情景。然后我们在给定的情景下,会对事件发生的概率进行估计。当各种情景(影响到事件的情况)之间是互斥和穷举的时候,那么就没有可能发生的结果被遗漏。于是,我们可以利用全概率公式(total probability rule)来分析事件发生的可能性。这条法则用以各种情景为条件的概率表达了事件发生的无条件概率。

下面的全概率公式对两种情形给出表述。第一部分给出的是一个最简单的情形,在这种情形下我们有两种情景出现。这里将引入一个新的记号:如果我们有一个事件或者情景 S ,那么不是 S 的事件,称为 S 的补集(complement),记为 S^c 。[⊖]注意 $P(S) + P(S^c) = 1$,因为 S 或者非 S 两者必须有一个会发生。第二部分对于 n 个互斥且穷举的事件或者情景的一般情形给出了该法则的表述。

• 全概率公式。

$$\begin{aligned} 1. P(A) &= P(AS) + P(AS^c) \\ &= P(A|S)P(S) + P(A|S^c)P(S^c) \end{aligned} \quad (4-5)$$

⊖ 对于熟悉概率数学处理方法的读者, S (一个经常用来表示样本空间概念的记号)适合用来表示情景。

$$\begin{aligned} 2. P(A) &= P(AS_1) + P(AS_2) + \cdots + P(AS_n) \\ &= P(A | S_1)P(S_1) + P(A | S_2)P(S_2) + \cdots + P(A | S_n)P(S_n) \end{aligned} \tag{4-6}$$

式中， $S_1, S_2, \dots, S_n$ 表示互斥且穷举的情景或者事件。

式 (4-6) 说明了：任何事件的概率 $P(A)$ 都可以表达成给定情景条件下事件发生概率（如 $P(A | S_1)$ 项）的一个加权平均值；其应用在条件概率上的权重为每种情景发生的概率（如 $P(S_1)$ 乘以 $P(A | S_1)$ 一项），其中的情景必须是互斥且穷举的。该法则还会在贝叶斯公式中用到，我们会在本章的后面加以讨论。

在下面的例子中，我们将使用全概率公式来形成对于 BankCorp 每股收益的一致观点。

例 4-7 BankCorp 每股收益 (2)

你继续考察你是否能预测 BankCorp 季度每股收益变化的方向。你定义了 4 个事件：

事件	概率
A = 下一个季度的前后两季度 EPS 的变化为正	0.55
A^c = 下一个季度的前后两季度 EPS 的变化为 0 或者为负	0.45
S = 前一个季度的前后两季度 EPS 的变化为正	0.55
S^c = 前一个季度的前后两季度 EPS 的变化为 0 或者为负	0.45

在了解这些数据之后，你可以观察到一定的 EPS 变化的持续性：增加之后倾向于继续增加，减少之后倾向于继续减少。你第一个估计出的概率为 $P(\text{下一个季度的前后两季度 EPS 的变化为正} | \text{前一个季度的前后两季度 EPS 的变化为 0 或者为负}) = P(A | S^c) = 0.40$ 。最近一个季度（2004 年二季度）的 EPS 已经公布了，该变化是一个正的前后变化（事件 S ）。你想要预测 2004 年第三季度的 EPS。

(1) 写出下述表述事件的概率记号：“给定前一个季度的前后两季度 EPS 的变化为正，下一个季度的前后两季度 EPS 的变化为正。”

(2) 计算出第 1 题的概率。（计算这个与其他的概率和信念一致的概率。）

解：

(1) 该表述事件的概率记号为 $P(A | S)$ 。

(2) 给定前一个季度的前后两季度 EPS 的变化为正，下一个季度的前后两季度 EPS 的变化为正的的概率为 0.673。具体计算如下。

根据式 (4-5)， $P(A) = P(A | S)P(S) + P(A | S^c)P(S^c)$ 。需要计算 $P(A | S)$ 的概率值已经都有了： $P(A) = 0.55$ ， $P(S) = 0.55$ ， $P(S^c) = 0.45$ 和 $P(A | S^c) = 0.40$ 。将这些值代入式 (4-5)，

$$0.55 = P(A | S)(0.55) + 0.40(0.45)$$

求解未知量， $P(A | S) = [0.55 - 0.40(0.45)] / 0.55 = 0.672727$ ，或者 0.673。

你得到 $P(\text{下一个季度的前后两季度 EPS 的变化为正} | \text{前一个季度的前后两季度 EPS 的变化为正}) = 0.673$ 。任何其他的概率与你估计的概率是不一致的。正的 EPS 变化的条件概率 0.673 比 EPS 增加的无条件概率 0.55 大，反映了 EPS 变化的持续性。

在统计学概念和市场收益率一章中，我们讨论了加权平均数或者加权均值的概念。在那一章

中所举的例子为投资组合收益率是投资组合中每个资产收益率的加权平均值，其中应用于每个资产收益率上的权重为投资组合中投资于该项资产的比例。全概率公式（它是用条件概率来表示无条件概率的一种法则）也是一种加权平均。在那个公式中，情景发生的概率被用来作为权重。加权平均定义中的一条就是权重之和等于1。而互斥且穷举事件的概率之和当然等于1（这是概率定义中的一部分）。我们接下去讨论的加权平均（随机变量的期望值）同样也用概率作为权重。

一个随机变量的期望值是在投资中最重要的定量概念。投资者一直在估计备选投资的收益率，预测EPS和其他公司财务变量和比率以及评估其他影响他们财务状况的因素时利用期望值。一个随机变量的期望值定义如下：

- **期望值的定义。**一个随机变量的期望值（expected value）是其可能结果的概率加权平均值。对于一个随机变量 X 来说， X 的期望值记为 $E(X)$ 。

期望值（例如，期望股票收益率）或者是针对未来，作为一个预测，或者是“真正的”均值（我们在统计学概念和市场收益率一章所讨论的总体均值）。我们应该区分期望值与历史均值或样本均值这些概念。样本均值也用一个数字来反映一个中心值。然而，样本均值是根据一个特定的观测集而计算出的中心值，它是对于那些观测值的等权重加权平均值。总的来说，两者的区别在于一个针对未来，一个针对历史，或者是一个针对总体，一个针对样本。

例 4-8 BankCorp 每股收益 (3)

继续我们对于 BankCorp 每股收益的分析。在表 4-3 中，你已经记录下了 BankCorp 当前财年每股收益的概率分布。

表 4-3 BankCorp 每股收益的概率分布

概率	每股收益	概率	每股收益
0.15	2.60 美元	0.16	2.00 美元
0.45	2.45 美元	1.00	
0.24	2.20 美元		

那么，当前财年 BankCorp 每股收益的期望值为多少？
根据期望值的定义，我们列出每个结果，根据其概率值进行加权，并将其每项进行加总。

$$\begin{aligned} E(\text{EPS}) &= 0.15(\$2.60) + 0.45(\$2.45) + 0.24(\$2.20) + 0.16(\$2.00) \\ &= \$2.3405 \end{aligned}$$

EPS 的期望值为 2.34 美元。

汇总在例 4-8 中计算期望值的一个等式如下：

$$E(X) = P(X_1)X_1 + P(X_2)X_2 + \cdots + P(X_n)X_n = \sum_{i=1}^n P(X_i)X_i \tag{4-7}$$

式中， X_i 表示随机变量 X 的 n 种可能结果之一。[⊖]

期望值是我们的预测，因为我们在讨论随机变量，我们不能指望某个预测能一定实现（虽然我们希望在平均意义上预测是准确的）。因此，度量我们所面临的风险是非常重要的。方差和标

⊖ 为了简单，在本章中我们将所有的随机变量都处理为离散型随机变量，即它们有可数个可能结果。对于连续型随机变量，我们将在常见概率分布一章中和离散型随机变量一起进行介绍，其相对于求和运算的对应算子是积分。

准差度量了结果偏离期望值或预测值的程度。

- **方差的定义。**一个随机变量的方差 (variance) 是偏离其期望值平方偏差的期望:

$$\sigma^2(X) = E\{[X - E(X)]^2\} \tag{4-8}$$

方差的两个记号为 $\sigma^2(X)$ 和 $\text{Var}(X)$ 。

方差是一个大于等于 0 的数, 因为它是平方项之和。如果方差为 0, 那么就没有偏差或者风险。结果是确定的, 故 X 不是随机的。方差大于 0 意味着结果存在偏离。在其他都不变的情况下, 方差的增大意味着离散度的增加。 X 的方差是一个以 X 单位的平方作为单位的量。例如, 如果随机变量是以百分比作为单位的收益率, 那么收益率的方差就是以百分比平方作为单位的量。标准差比方差易于解释, 因为它和随机变量的单位是相同的。如果随机变量是以百分比作为单位的收益率, 那么收益率的标准差也是以百分比作为单位的。

- **标准差的定义。**标准差 (standard deviation) 是方差正的平方根。

熟悉这些概念的最好方法就是做一些具体的例子。

例 4-9 BankCorp 每股收益 (4)

在例 4-8 中, 已经计算出 BankCorp 每股收益的期望值为 2.34 美元。这就是你的预测值。现在你想要度量结果可能偏离预测值的程度。表 4-4 给出了你对于当前财年每股收益的概率分布的看法。

表 4-4 BankCorp 每股收益的概率分布

概率	每股收益
0.15	2.60 美元
0.45	2.45 美元
0.24	2.20 美元
0.16	2.00 美元
1.00	

那么当前财年 BankCorp 每股收益的方差和标准差为多少呢?

计算顺序总是为先求期望值, 再求方差, 最后求标准差。期望值之前已经算出来了。接着根据之前方差的定义, 计算每个结果偏离均值或者期望值的偏差, 然后将每个偏差平方, 并以其发生的概率对于每个平方偏差进行加权 (相乘), 最后将这些项加总。

$$\begin{aligned} \sigma^2(\text{EPS}) &= P(\$2.60)[\$2.60 - E(\text{EPS})]^2 + P(\$2.45)[\$2.45 - E(\text{EPS})]^2 \\ &\quad + P(\$2.20)[\$2.20 - E(\text{EPS})]^2 + P(\$2.00)[\$2.00 - E(\text{EPS})]^2 \\ &= 0.15[2.60 - 2.34]^2 + 0.45[2.45 - 2.34]^2 + 0.24[2.20 - 2.34]^2 + 0.16[2.00 - 2.34]^2 \\ &= 0.01014 + 0.005445 + 0.004704 + 0.018496 = \$0.038785 \end{aligned}$$

标准差是 \$0.038785 正的平方根:

$$\sigma(\text{EPS}) = \$0.038785^{1/2} = \$0.196939, \text{ 或约为 } 0.20 \text{ 美元。}$$

汇总在例 4-9 中计算方差的等式:

$$\begin{aligned} \sigma^2(X) &= P(X_1)[X_1 - E(X)]^2 + P(X_2)[X_2 - E(X)]^2 + \cdots + P(X_n)[X_n - E(X)]^2 \\ &= \sum_{i=1}^n P(X_i)[X_i - E(X)]^2 \end{aligned} \tag{4-9}$$

式中, X_i 表示随机变量 X 的 n 种可能结果之一。

在投资过程中, 我们利用任何可获得的相关信息来帮助我们作出预测。当我们作出我们的期望或者预测时, 我们往往会基于新信息或者新的事件作出调整; 在这些情况下我们使用条件期望

值 (conditional expected values)。一个随机变量 X 在给定事件或者情景 S 下的期望值记为 $E(X | S)$ 。假如随机变量 X 能取到 n 个不同结果 $X_1, X_2, \dots, X_n$ 中的任何一个 (这些结果形成一个互斥且穷举的集合)。基于条件 S 的 X 的期望值为第 1 个结果 X_1 乘以在 S 条件下第 1 个结果出现的概率 $P(X_1 | S)$, 加上第 2 个结果 X_2 乘以在 S 条件下第 2 个结果出现的概率 $P(X_2 | S)$, 以此类推。

$$E(X | S) = P(X_1 | S)X_1 + P(X_2 | S)X_2 + \dots + P(X_n | S)X_n \tag{4-10}$$

我们将简要说明一下这个等式。

与将无条件概率表示为条件概率的全概率公式一样, 我们也有将无条件期望表示为条件期望值的法则。该法则称为对于期望值的全概率公式 (total probability rule for expected value)。

• 对于期望值的全概率公式。

1. $E(X) = E(X | S)P(S) + E(X | S^c)P(S^c)$ (4-11)

2. $E(X) = E(X | S_1)P(S_1) + E(X | S_2)P(S_2) + \dots + E(X | S_n)P(S_n)$ (4-12)

式中, $S_1, S_2, \dots, S_n$ 表示互斥且穷举的情景或者事件。

一般的情形 (式 (4-12)) 说明 X 的期望值等于给定情景 1 条件下 X 的期望值 $E(X | S_1)$ 乘以情景 1 的概率 $P(S_1)$, 加上给定情景 2 条件下 X 的期望值 $E(X | S_2)$ 乘以情景 2 的概率 $P(S_2)$, 以此类推。

为了使用该法则, 我们要设定出对于理解随机变量结果有益的互斥且穷举的情景。如同我们现在要讨论的, 该方法在例 4-8 和例 4-9 建立 BankCorp 每股收益的概率分布中已经用到。

BankCorp 的收益对于利率的变化是敏感的, 下降的利率环境对其收益有利。假如有 0.60 的可能性, BankCorp 将会在当前财年面临一个利率下降的经营环境, 而有 0.40 的可能性, 它会面临一个稳定不变的利率环境 (所估计的利率上升环境的可能性可以忽略不计)。如果利率下降的环境出现, EPS 将为 2.60 美元的概率估计为 0.25, 而 EPS 为 2.45 美元的概率估计为 0.75。注意利率下降环境出现的概率 0.60 乘以给定利率下降环境 EPS 为 2.60 美元的概率 0.25 等于上面例 4-8 和例 4-9 表中 EPS 为 2.60 美元的 (无条件) 概率 0.15。这个概率是一致的。同样, $0.60 (0.75) = 0.45$, 等于表 4-3 和表 4-4 中给出的 EPS 为 2.45 美元的概率。图 4-2 中的树状图 (tree diagram) 给出了余下的分析。

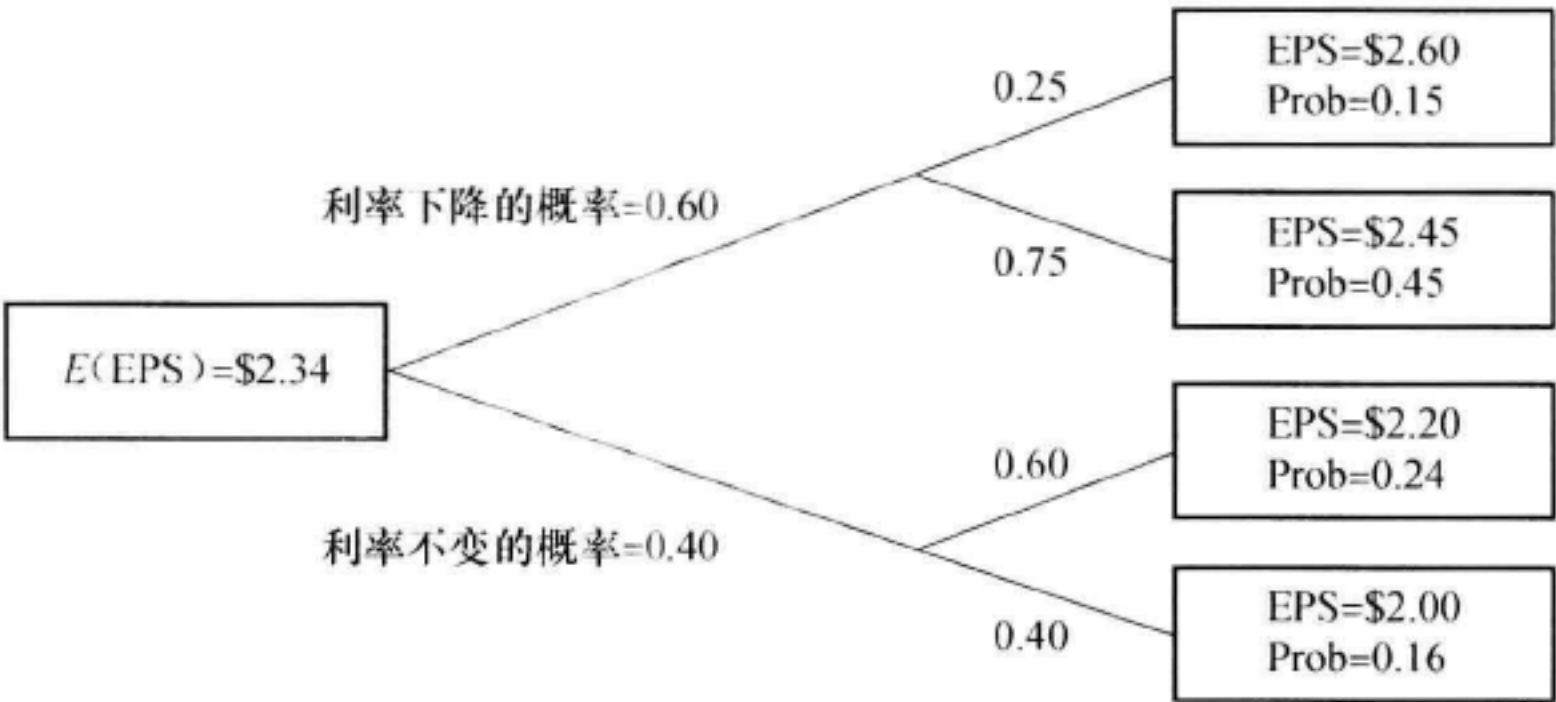


图 4-2 BankCorp 预测的 EPS

一个利率下降的环境将我们指向树的一个结点, 该结点上有两个分叉, 分别指向 2.60 美元和 2.45 美元这两个结果。利用式 (4-10), 我们能求出给定利率下降环境时的期望 EPS。具体如下:

$$E(\text{EPS} | \text{利率下降环境}) = 0.25(\$2.60) + 0.75(\$2.45) = \$2.4875$$

如果利率不变,

$$E(\text{EPS} \mid \text{利率不变环境}) = 0.60(\$2.20) + 0.40(\$2.00) = \$2.12$$

一旦我们有利率不变的新信息,我们将我们原先 EPS 的期望值从 2.34 美元向下修改至 2.12 美元。现在利用对于期望值的全概率公式,

$$E(\text{EPS}) = E(\text{EPS} \mid \text{利率下降环境})P(\text{利率下降环境}) + E(\text{EPS} \mid \text{利率不变环境})P(\text{利率不变环境})$$

$$\text{所以, } E(\text{EPS}) = \$2.4875(0.60) + \$2.12(0.40) = \$2.3405, \text{ 或者约 } 2.34 \text{ 美元。}$$

这个量与例 4-8 中直接通过概率分布计算出的 EPS 期望值的估计值相同。正如我们的概率一定是一致的,我们的期望值也必须是一致的,而无论是无条件的还是有条件的;不然我们的投资行为将可能使其他投资者以我们的成本获得收益机会。

回顾一下,我们首先建立一些影响所讨论事件发生结果的因素或者情景。在分配给这些情景概率之后,我们形成基于这些不同情景的条件期望。于是,我们可以倒推出今天的期望值。在之前的问题中, EPS 是所要考虑的事件,而利率环境是影响 EPS 的因素。

我们也可以在给定每个情景条件下计算 EPS 的方差。

$$\begin{aligned}\sigma^2(\text{EPS} \mid \text{利率下降环境}) &= P(\$2.60 \mid \text{利率下降环境}) \times [\$2.60 - E(\$2.60 \mid \text{利率下降环境})]^2 \\ &\quad + P(\$2.45 \mid \text{利率下降环境}) [\$2.45 - E(\$2.45 \mid \text{利率下降环境})]^2 \\ &= 0.25(\$2.60 - \$2.4875)^2 + 0.75(\$2.45 - \$2.4875)^2 \\ &= 0.004219\end{aligned}$$

$$\begin{aligned}\sigma^2(\text{EPS} \mid \text{利率不变环境}) &= P(\$2.20 \mid \text{利率不变环境}) [\$2.20 - E(\$2.20 \mid \text{利率下降环境})]^2 \\ &\quad + P(\$2.00 \mid \text{利率不变环境}) \times [\$2.00 - E(\$2.00 \mid \text{利率不变环境})]^2 \\ &= 0.60(\$2.20 - \$2.12)^2 + 0.40(\$2.00 - \$2.12)^2 \\ &= 0.0096\end{aligned}$$

给定利率下降环境下的 EPS 方差和给定利率不变环境下的 EPS 方差,这些方差是条件方差 (conditional variances)。无条件方差和条件方差间的关系是一个相对高级的讨论主题。[⊖]主要的要点为:
①方差像期望值一样,具有相对于无条件方差概念的条件方差的概念。②在给定一定情景条件下,我们可以利用条件概率来评估风险。

例 4-10 BankCorp 每股收益 (5)

我们继续讨论 BankCorp,现在我们关注 BankCorp 的成本结构。一个你用来研究 BankCorp 运营成本的模型为:

$$\hat{Y} = a + bX$$

式中, $\hat{Y}$ 表示以百万美元计的运营成本的预测值,而 X 表示分支营业处的数目。 $\hat{Y}$ 表示给定 X 条件下 Y 的期望值,或者为 $E(Y \mid X)$ 。($\hat{Y}$ 是在回归分析中所用的记号,我们将在后面的章节中对其进行讨论。)你将截距 a 解释为固定成本,而 b 解释为可变成本。你估计出的等式为:

$$\hat{Y} = 12.5 + 0.65X$$

⊖ EPS 的无条件方差是下面两项之和:(1) 条件方差的期望值 (概率加权平均)(与全概率公式相同);(2) EPS 条件期望值的方差。第 2 项的产生是因为条件期望值的变动是风险的一个来源。(在例子中) 第 1 项为 $\sigma^2(\text{EPS}) = P(\text{利率下降环境}) \sigma^2(\text{EPS} \mid \text{利率下降环境}) + P(\text{利率不变环境}) \sigma^2(\text{EPS} \mid \text{利率不变环境}) = 0.60(0.004219) + 0.40(0.0096) = 0.006371$ 。第 2 项为 $\sigma^2[E(\text{EPS} \mid \text{利率环境})] = 0.60(\$2.4875 - \$2.34)^2 + 0.40(\$2.12 - \$2.34)^2 = 0.032414$ 。将上面两项相加,无条件方差为 $0.006371 + 0.032414 = 0.038785$ 。

BankCorp 现在有 66 个营业处，故该等式所估计的运营成本为 $12.5 + 0.65(66) = \$55.4$ (百万)。你有两种增长的情景，具体列示在图 4-3 中的树状图中。

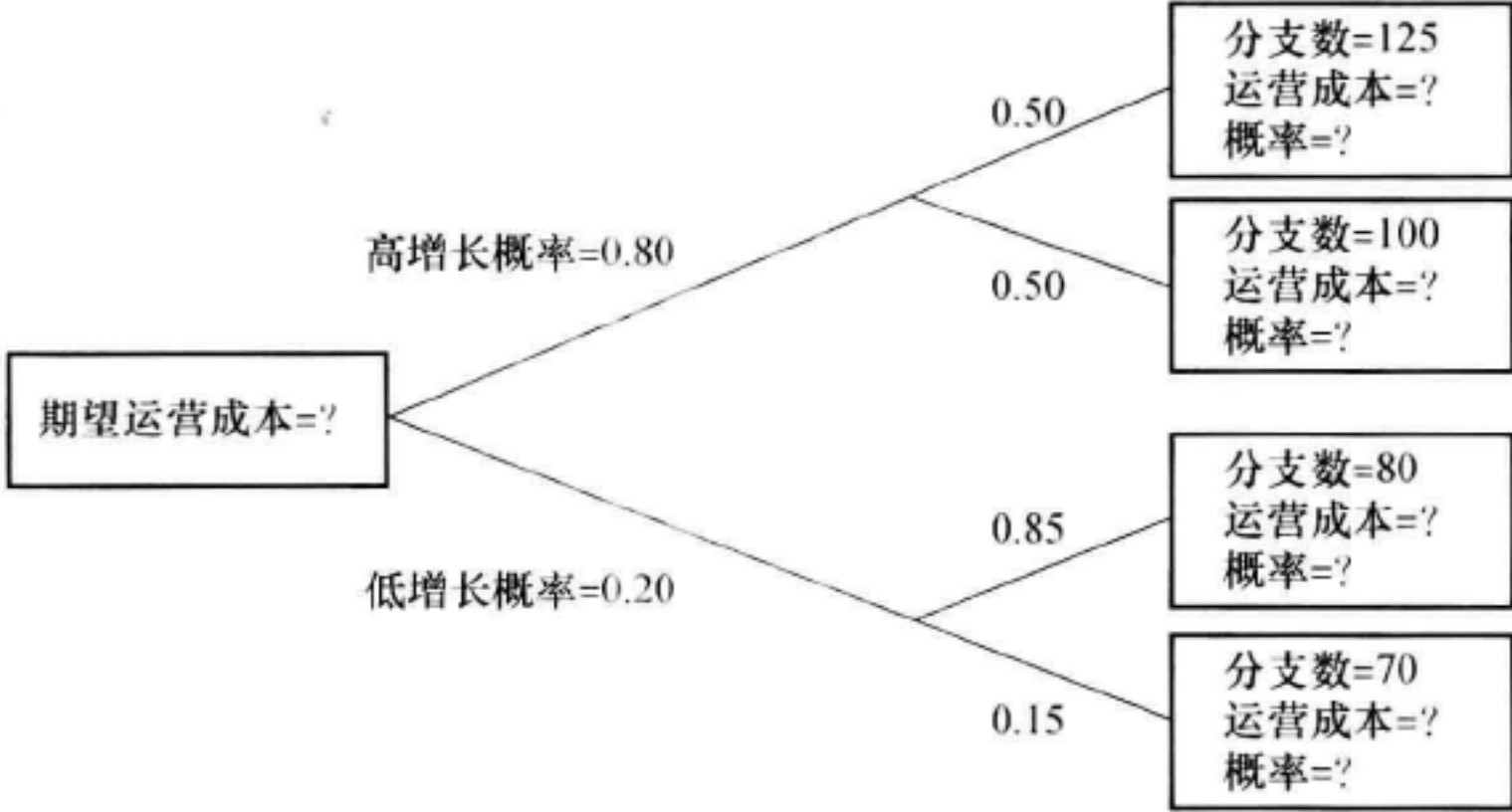


图 4-3 BankCorp 预测的运营成本

- (1) 利用 $\hat{Y} = a + bX$ ，计算给定不同运营成本水平下的预测运营成本。指出每个分支营业处各个水平的概率。将本问题的这些答案填在树状图最后那些空格中。
- (2) 计算在高增长情景下运营成本的预测值，同时也计算一下在低增长情景下运营成本的预测值。
- (3) 在树状图最初的空格中回答下面的问题：BankCorp 的期望运营成本是多少？

解：

(1) 利用 $\hat{Y} = 12.5 + 0.65X$ ，从上至下，我们有：

运营成本	概率
$\hat{Y} = 12.5 + 0.65(125) = \93.75 (百万)	$0.80(0.50) = 0.40$
$\hat{Y} = 12.5 + 0.65(100) = \77.50 (百万)	$0.80(0.50) = 0.40$
$\hat{Y} = 12.5 + 0.65(80) = \64.50 (百万)	$0.20(0.85) = 0.17$
$\hat{Y} = 12.5 + 0.65(70) = \58.00 (百万)	$0.20(0.15) = 0.03$
$\hat{Y} = 12.5 + 0.65(125) = \93.75 (百万)	总和 = 1.00

(2) 以百万计的美元数量为：

$E(\text{运营成本} \mid \text{高增长}) = 0.50(\$93.75) + 0.50(\$77.50) = \85.625

$E(\text{运营成本} \mid \text{低增长}) = 0.85(\$64.50) + 0.15(\$58.00) = \63.525

(3) 以百万计的美元数量为：

$$\begin{aligned} E(\text{运营成本}) &= E(\text{运营成本} \mid \text{高增长})P(\text{高增长}) + E(\text{运营成本} \mid \text{低增长})P(\text{低增长}) \\ &= \$85.625(0.80) + \$63.525(0.20) \\ &= \$81.205 \end{aligned}$$

BankCorp 的预期运营成本为 81.205 百万美元。

我们将在讨论贝叶斯公式的时候再次看到条件概率。这一节已经介绍了许多可以利用概率概念表述的问题。下面这个问题将在利用这些概念的同时，再介绍一些分析的技巧。

例 4-11 单期债务工具的违约风险溢价

作为一个短期债券投资组合的合伙经理，你正在考察一个投机级别、一年期无息债券的定价。对于这一类型的债券，其收益为所支付的金额和到期所收到的本金之差。你的目的是要估计该债券的一个适当的违约风险溢价。我们将违约风险溢价定义为高于无风险收益率的额外收益率，该收益率被用来补偿投资者所面临的违约风险。如果 R 是该债务工具承诺的收益率（到期收益率）， R_f 是无风险收益率，那么违约风险溢价为 $R - R_f$ 。你估计债券违约的风险为 P （债券违约）= 0.60。通过观察当前货币市场的收益率，你发现一年期美国国库券提供 5.8% 的收益率，其可以作为 R_f 的一个估计。第 1 步，你要做出一个简化的假设：一旦债券违约，债券持有者将不能获得任何补偿。那么（在此条件下）你投资该工具所需要的最低违约风险溢价应该为多少？

这类问题的难点在于找到一个考虑的出发点。在许多问题中（包括这个问题），第 1 步较好的选择是将所有可能结果按照经济上的逻辑，分成互斥其穷举的一组事件。这里，从债券持有者的角度来看，两个事件会影响到收益率，即债券违约和债券不违约。这两个事件涵盖了所有可能的结果，那么这两个事件如何影响债券持有者的收益率呢？第 2 步就要计算在这两个事件下的债券价值。我们没有对债券的面值有具体的设定，但是我们可以计算每 1 美元的价值或者一单位投资货币的价值。

	债券违约	债券不违约
债券价值	\$0	$\$(1 + R)$

第 3 步是计算债券的期望值（每 1 美元投资的值）。

$$E(\text{债券}) = \$0 \times P(\text{债券违约}) + \$(1 + R)[1 - P(\text{债券违约})]$$

因此， $E(\text{债券}) = \$(1 + R)[1 - P(\text{债券违约})]$ 。每 1 美元投资的国库券的期望值为 $(1 + R_f)$ 。事实上，这个值是确定的，因为国库券是无风险的。下面一步需要经济学的推理。你想要违约风险溢价充分大以使得你期望至少在盈亏平衡点上你能获得与国库券相同的收益率。如果每 1 美元投资的债券的期望价值与国库券的期望价值相等，那么这一结果将会发生。

$$\text{债券的期望价值} = \text{国库券的期望价值}$$

$$\$(1 + R)[1 - P(\text{债券违约})] = (1 + R_f)$$

求解债券承诺的收益率，你会得到： $R = \{(1 + R_f) / [1 - P(\text{债券违约})]\} - 1$ 。用题中给出的值代入其中，有 $R = [1.058 / (1 - 0.06)] - 1 = 1.12553 - 1 = 0.12553$ 或者约 12.55%，违约风险溢价为 $R - R_f = 12.55\% - 5.8\% = 6.75\%$ 。

你需要的违约风险溢价至少为 675 个基点。你可以将该事实陈述如下：如果债券的定价为 12.55%，那么你将以 94% 的可能性获得一个 675 基点的价差以及债券的本金。然而，如果债券违约，你将一无所获。在 675 个基点溢价的情况下，你预计该投资相对于投资于国库券正好盈亏相抵，因为一项无息债券的投资收益存在可变性，如果你是风险厌恶型的，你将需要一个比 675 个基点更大的溢价。

这个分析是一个开端。债券持有者经常在违约之后获得部分投资的补偿。下一步将把一个补偿率考虑到该问题中去。

在这一节，我们已经将诸如 EPS 等随机变量作为单独的量进行了处理，但我们还没有讨论如何对于作为其他随机变量函数 EPS 的期望值和方差进行描述。投资组合收益率明显就是作为其他随机变量（投资组合中每个证券随机收益率）函数的一个随机变量。为了分析一个投资组合的期望收益率及其方差，我们必须理解到这些量是单个证券收益率的函数。从投资组合收益率的离散程度或者方差中，我们看到单个证券收益率间共同变化或者共变的方式是非常重要的。为了理解这些变化的重要性，我们需要探讨一些新的概念——协方差和相关性。下一节就会在投资组合期望收益率及其方差的背景下，介绍这些概念。

4.3 投资组合的期望收益和收益的方差

现代投资组合理论经常利用这样的观点：把期望收益率作为回报的一个度量，而把收益率的方差作为风险的一个度量，以此来对于投资机会进行评估。对于投资组合期望收益率及其方差的计算和解释是一项基本的技能。在这一节中，我们将建立投资组合期望收益率和收益率的方差的概念。^①投资组合收益率是由每个持有证券的收益率所决定的。而单个资产收益率函数的投资组合方差的计算却比前一节中介绍的方差计算要复杂一些。

我们以下面的这个投资组合作为一个例子。该投资组合中 50% 投资于一个标准普尔 500 指数基金，25% 投资于一个美国长期公司债券基金，剩下 25% 投资于一个以 MSCI EAFE 指数（摩根士丹利欧澳远东指数，代表欧洲、澳洲以及远东股票市场）作为跟踪标的基金。表 4-5 列出了这些权重。

表 4-5 投资组合权重

资产类	权重
S&P500	0.50
美国长期公司债券	0.25
MSCI EAFE	0.25

我们首先讨论投资组合期望收益率的计算方法。在前一节中，我们将一个随机变量的期望值定义为以其可能结果发生概率的加权平均值。而我们知道，投资组合收益率是投资组合中每个证券收益率的加权平均值。同样的，投资组合的期望收益率就是利用相同权重的，投资组合中每个证券期望收益率的加权平均值。当我们已经估计出每个证券的期望收益率时，我们马上可以获得投资组合的期望收益率。这一简单的事实满足下面期望值的这些性质。

• **期望值的性质。**令 w_i 为任一常数， R_i 是一个随机变量。

1. 一个常数乘以一个随机变量的期望值等于该常数乘以该随机变量的期望值。

$$E(w_i R_i) = w_i E(R_i)$$

2. 随机变量加权之和的期望值等于利用相同权重对于每个随机变量期望值的加权之和。

$$E(w_1 R_1 + w_2 R_2 + \cdots + w_n R_n) = w_1 E(R_1) + w_2 E(R_2) + \cdots + w_n E(R_n) \tag{4-13}$$

假设我们有一个给定期望值的随机变量。如果我们在其每个可能结果上乘以 2，那么该随机变量的期望值也会变为原来的两倍，这是第一个性质的含义。第 2 个性质是直接产生投资组合期望收益率表达式的一个法则。一个有 n 个资产的投资组合，其权重被定义为 $w_1, w_2, \cdots, w_n$ ，并且权重之和为 1，故投资组合收益率 R_p 为 $R_p = w_1 R_1 + w_2 R_2 + \cdots + w_n R_n$ 。于是，我们可以得到下面的原理：

① 虽然我们在这一节中列出了一些基本概念，但是我们没有列出均值方差分析本身。对于均值方差分析的表述，读者可以阅读投资组合概念一章以及如 Bodie, Kane 和 Marcus (2001), Elton, Gruber, Brown 和 Goetzmann (2003), Reilly 和 Brown (2003), Sharpe, Alexander 和 Bailey (1999) 的标准教科书中的深入讨论。

- 投资组合期望收益率的计算方法。给定一个 n 个资产的投资组合，其预期收益率等于该投资组合中每个证券期望收益率的加权平均值：

$$E(R_p) = E(w_1R_1 + w_2R_2 + \cdots + w_nR_n) = w_1E(R_1) + w_2E(R_2) + \cdots + w_nE(R_n)$$

假设我们已经得到了投资组合中每个资产期望收益率的估计值，如表 4-6 所给出的。我们就可以计算出投资组合的期望收益率为 11.75%：

$$\begin{aligned} E(R_p) &= w_1E(R_1) + w_2E(R_2) + w_3E(R_3) \\ &= 0.50(13\%) + 0.25(6\%) + 0.25(15\%) \\ &= 11.75\% \end{aligned}$$

表 4-6 权重和期望收益率

资产类	权重	期望收益率 (%)
标准普尔 500	0.50	13
美国长期公司债券	0.25	6
MSCI EAFE	0.25	15

在前一节中，我们研究了作为度量出现结

果偏离期望值的偏差程度的方差。现在我们要关注的是作为度量投资风险的投资组合收益率的方差。令 R_p 是投资组合的收益率，根据式 (4-8)，投资组合的方差为 $\sigma^2(R_p) = E\{[R_p - E(R_p)]^2\}$ 。那么我们如何使用这个定义呢？在统计学概念和市场收益率一章中，我们曾学习了如何基于样本收益率来计算历史的或样本的方差，现在我们要在前瞻性的意义下考察方差。我们将使用投资组合中每个资产的信息来获得投资组合收益率的方差。为了避免记号的混乱，我们将 $E(R_p)$ 简记为 ER_p 。另外，我们需要协方差的概念。

- 协方差的定义。给定两个随机变量 R_i 和 R_j ， R_i 和 R_j 的协方差为：

$$\text{Cov}(R_i, R_j) = E[(R_i - ER_i)(R_j - ER_j)] \quad (4-14)$$

协方差的另一个记法为 $\sigma(R_i, R_j)$ 和 σ_{ij} 。

式 (4-14) 表示，两个随机变量的协方差是每个随机变量偏离其期望值相互乘积的概率加权平均值。我们将在建立所需概念之后，再回过头来对于协方差进行讨论。

从方差的定义中，我们可以得到：

$$\begin{aligned} \sigma^2(R_p) &= E[(R_p - ER_p)^2] \\ &= E\{[w_1R_1 + w_2R_2 + w_3R_3 - E(w_1R_1 + w_2R_2 + w_3R_3)]^2\} \\ &= E\{[w_1R_1 + w_2R_2 + w_3R_3 - w_1ER_1 - w_2ER_2 - w_3ER_3]^2\} \quad (\text{利用式(4-13)}) \\ &= E\{[w_1(R_1 - ER_1) + w_2(R_2 - ER_2) + w_3(R_3 - ER_3)]^2\} \quad (\text{整理}) \\ &= E\{[w_1(R_1 - ER_1) + w_2(R_2 - ER_2) + w_3(R_3 - ER_3)] \\ &\quad \times [w_1(R_1 - ER_1) + w_2(R_2 - ER_2) + w_3(R_3 - ER_3)]\} \quad (\text{平方的含义}) \\ &= E[w_1w_1(R_1 - ER_1)(R_1 - ER_1) + w_1w_2(R_1 - ER_1)(R_2 - ER_2) \\ &\quad + w_1w_3(R_1 - ER_1)(R_3 - ER_3) + w_2w_1(R_2 - ER_2)(R_1 - ER_1) \\ &\quad + w_2w_2(R_2 - ER_2)(R_2 - ER_2) + w_2w_3(R_2 - ER_2)(R_3 - ER_3) \\ &\quad + w_3w_1(R_3 - ER_3)(R_1 - ER_1) + w_3w_2(R_3 - ER_3)(R_2 - ER_2) \\ &\quad + w_3w_3(R_3 - ER_3)(R_3 - ER_3)] \quad (\text{进行乘法计算}) \\ &= w_1^2E[(R_1 - ER_1)^2] + w_1w_2E[(R_1 - ER_1)(R_2 - ER_2)] \\ &\quad + w_1w_3E[(R_1 - ER_1)(R_3 - ER_3)] + w_2w_1E[(R_2 - ER_2)(R_1 - ER_1)] \\ &\quad + w_2^2E[(R_2 - ER_2)^2] + w_2w_3E[(R_2 - ER_2)(R_3 - ER_3)] \\ &\quad + w_3w_1E[(R_3 - ER_3)(R_1 - ER_1)] + w_3w_2E[(R_3 - ER_3)(R_2 - ER_2)] \end{aligned}$$

$$\begin{aligned}
 &+ w_3^2 E[(R_3 - ER_3)^2] \quad (\text{回忆一下, } w_i \text{ 项是常数}) \\
 &= w_1^2 \sigma^2(R_1) + w_1 w_2 \text{Cov}(R_1, R_2) + w_1 w_3 \text{Cov}(R_1, R_3) \\
 &\quad + w_1 w_2 \text{Cov}(R_1, R_2) + w_2^2 \sigma^2(R_2) + w_2 w_3 \text{Cov}(R_2, R_3) \\
 &\quad + w_1 w_3 \text{Cov}(R_1, R_3) + w_2 w_3 \text{Cov}(R_2, R_3) + w_3^2 \sigma^2(R_3) \quad (4-15)
 \end{aligned}$$

最后一步依据的是方差和协方差的定义。[⊖]对于式(4-15)中斜体的协方差项,我们利用这样一个事实:协方差中变量的先后顺序对于结果是没有影响的:例如 $\text{Cov}(R_i, R_j) = \text{Cov}(R_j, R_i)$ 。正如我们将要展现出来的,对角线上的方差项 $\sigma^2(R_1)$ 、 $\sigma^2(R_2)$ 和 $\sigma^2(R_3)$ 也可以分别表示为 $\text{Cov}(R_1, R_1)$ 、 $\text{Cov}(R_2, R_2)$ 和 $\text{Cov}(R_3, R_3)$ 。利用这一事实,最简约的表达式(4-15)的方式为 $\sigma^2(R_p) = \sum_{i=1}^3 \sum_{j=1}^3 w_i w_j \text{Cov}(R_i, R_j)$ 。两重的求和符号的含义是:“先设定 $i=1$ 将 j 从 1 变到 3;然后设定 $i=2$ 让 j 从 1 变到 3;再设定 $i=3$ 将 j 从 1 变到 3;最后将这九项相加。”对于某个规模为 n 的资产组合的一般表述为:

$$\sigma^2(R_p) = \sum_{i=1}^n \sum_{j=1}^n w_i w_j \text{Cov}(R_i, R_j) \quad (4-16)$$

我们从式(4-15)中看到每个资产收益率的方差(对角线上粗体的项)只是组成投资组合方差的一部分,而不是全部。这3个方差实际上在数目上要少于不在对角线上的协方差项的6项。对于3个资产来说,此两者的比率为1:2,或者50%。如果是20个资产,那么就会有20个方差项和 $20(20) - 20 = 380$ 个不在对角线上的协方差项。方差项和不在对角线上的协方差项的比率会小于6:100,或者6%。于是,我们首先观察到的就是在其他条件不变的情况下,随着持有资产的增加,协方差[⊖]在投资组合方差中所起到的作用将会越来越大。

那么,协方差对于投资组合方差的真正影响是什么呢?协方差项描述了收益率间的变化是如何对于投资组合方差产生影响的。例如,考虑两只股票:当一个获得低收益率时(相对于其期望收益率来说),另一个倾向于获得高收益率(相对于其期望收益率)。于是,一只股票的收益率倾向于抵消另一只股票的收益率,这降低了整个投资组合收益率的变动性或者其方差。与方差一样,协方差的单位也难以进行解释。因此,我们将引入一个更加直观的概念。同时,从协方差的定义中,我们能形成两个关于协方差的观测结果。

1. 我们可以将协方差的符号解释如下:

当一个资产的收益率高于其期望值,而另一个资产的收益率倾向于低于其期望值时(收益率间平均意义下的反向关系),收益率的协方差为负。

当资产的收益率间不相关时,协方差为0。

当两个资产的收益率倾向于同时位于它们期望值的同一侧(同时大于或者小于)时(在平均意义下的收益率间的关系),收益率的协方差为正。

2. 一个随机变量和其自身的协方差(自己和自己的协方差)就是其自身的方差:

$$\text{Cov}(R, R) = E\{[R - E(R)][R - E(R)]\} = E\{[R - E(R)]^2\} = \sigma^2(R)。$$

⊖ 方差和协方差的一些有用的结论包括:①一个常数乘以一个随机变量的方差等于该常数的平方乘以该随机变量的方差,或者 $\sigma^2(wR) = w^2 \sigma^2(R)$;②一个常数加上一个随机变量的方差等于该随机变量的方差,或者 $\sigma^2(w + R) = \sigma^2(R)$, 因为一个常数的方差为0。③一个常数和随机变量的协方差为0。

⊖ 当非对角线(不在对角线上的)协方差作为协方差的含义明显的时候,就如这里所显示的,我们将省略其修饰语。在这个意义下,协方差一词经常被使用。

一个完整的协方差列表需要用到用于计算投资组合方差的所有统计数据。协方差经常用一个方阵来表示，其被称为协方差矩阵（covariance matrix）。表4-7 汇总了投资组合期望收益率及其方差的输入变量。

表 4-7 投资组合期望收益率及其方差的输入变量

A. 投资组合期望收益率的输入变量			
资产	A	B	C
	$E(R_A)$	$E(R_B)$	$E(R_C)$
B. 协方差矩阵：投资组合收益率方差的输入变量			
资产	A	B	C
A	Cov(R_A, R_A)	Cov(R_A, R_B)	Cov(R_A, R_C)
B	Cov(R_B, R_A)	Cov(R_B, R_B)	Cov(R_B, R_C)
C	Cov(R_C, R_A)	Cov(R_C, R_B)	Cov(R_C, R_C)

在 3 个资产的情况下，协方差矩阵有 $3^2 = 3 \times 3 = 9$ 个元素，但是习惯上，我们将对角线上的方差项与非对角线的协方差项分开处理。在表 4-7 中的对角线项被以粗体表示。这一区分是自然的，因为证券的方差是单个变量的概念。因此，排除方差，我们有 $9 - 3 = 6$ 个协方差。但是由于 $\text{Cov}(R_B, R_A) = \text{Cov}(R_A, R_B)$ ， $\text{Cov}(R_C, R_A) = \text{Cov}(R_A, R_C)$ 以及 $\text{Cov}(R_C, R_B) = \text{Cov}(R_B, R_C)$ ，协方差矩阵中在对角线下方的元素是在对角线上方元素的镜像。因此，我们只有 $6/2 = 3$ 个不同的协方差项需要估计。一般地，对于 n 个证券，我们有 $n(n - 1)/2$ 个不同的协方差和 n 个方差项需要估计。

假设我们有如表 4-8 所表示的协方差矩阵。利用式 (4-15)，并将方差项放在一起，就有：

$$\begin{aligned}\sigma^2(R_p) &= w_1^2\sigma^2(R_1) + w_2^2\sigma^2(R_2) + w_3^2\sigma^2(R_3) + 2w_1w_2\text{Cov}(R_1,R_2) \\ &\quad + 2w_1w_3\text{Cov}(R_1,R_3) + 2w_2w_3\text{Cov}(R_2,R_3) \\ &= (0.50)^2(400) + (0.25)^2(81) + (0.25)^2(441) + 2(0.50)(0.25)(45) \\ &\quad + 2(0.50)(0.25)(189) + 2(0.50)(0.25)(38) \\ &= 100 + 5.0625 + 27.5625 + 11.25 + 47.25 + 4.75 = 195.875\end{aligned}$$

(4-17)

表 4-8 协方差矩阵

	标准普尔 500	美国长期公司债券	MSCI EAFE
标准普尔 500	400	45	189
美国长期公司债券	45	81	38
MSCI EAFE	189	38	441

方差为 195.875，收益率的标准差为 $195.875^{1/2} = 14\%$ ，最后得到投资组合具有 11.75% 的期望年收益率和 14% 的收益率标准差。

让我们看一下上面计算中的 3 项，它们的和 $100 + 5.0625 + 27.5625 = 132.625$ 是每个资产方差对于投资组合方差的贡献。如果这 3 个资产的收益率间是相互独立的，那么协方差将为 0，且投资组合收益率的标准差，相对于之前的 14%，为 $132.625^{1/2} = 11.52\%$ 。投资组合将具有更小的风险。假如协方差项是负的。那么一个负数将会加到 132.625 之上，因此投资组合的方差和风险将更加小。与此同时，我们没有改变期望收益率。故对于相同的投资组合期望收益率，该投资组合将具有更小的风险。这一风险的降低是由分散化所带来的好处，即从持有一个多种资产的投资组合所带来的风险降低的好处。分散化的好处会随着协方差的降低而增加。这一观察是现代投资

组合理论的一个重要结果。当我们利用相关性的概念时，该结果可以表达得更加直观。于是我们可以说只要证券收益率不是完全正相关的，就能获得分散化的好处。更进一步，当其他条件一样时，证券收益率间的相关性越小，不分散的成本就越大。

- **相关系数的定义。**两个随机变量 R_i 和 R_j 的相关系数 (correlation) 定义为 $\rho(R_i, R_j) = \text{Cov}(R_i, R_j) / \sigma(R_i)\sigma(R_j)$ 。另外一个记号为 $\text{Corr}(R_i, R_j)$ 或者 ρ_{ij} 。

通常，我们会利用关系 $\text{Cov}(R_i, R_j) = \rho(R_i, R_j)\sigma(R_i)\sigma(R_j)$ 来替代协方差的表示。在定义中的除法使得相关系数称为一个纯数（没有度量单位的一个数），并且使其处于一个有最大值和最小值的区间之中。利用上述的定义，我们能只以协方差矩阵中的数据获得相关系数矩阵。表 4-9 列出了相关系数矩阵。

表 4-9 相关系数矩阵

	标准普尔 500	美国长期公司债券	MSCI EAFE
标准普尔 500	1.00	0.25	0.45
美国长期公司债券	0.25	1.00	0.20
MSCI EAFE	0.45	0.20	1.00

例如，在表 4-8 中，长期债券和 MSCI EAFE 的协方差为 38。从表 4-8 中对角线上的项得到，长期债券收益率的标准差为 $81^{1/2} = 9\%$ ，而 MSCI EAFE 收益率的标准差为 $441^{1/2} = 21\%$ 。相关系数 ρ （长期债券收益率，EAFE 收益率）为 $38 / (9\%)(21\%) = 0.201$ ，近似为 0.20。标准普尔 500 和自身的相关系数等于 1：其计算为标准普尔 500 与自己的协方差除以其自身标准差的平方。

- **相关系数的性质。**
 1. 两个随机变量 X 和 Y 的相关系数是一个位于 -1 和 $+1$ 之间的数：
$$-1 \leq \rho(X, Y) \leq +1$$
 2. 相关系数为 0（不相关变量）表示两个变量间没有任何线性（直线）关系。[⊖] 正相关关系的增加意味着正的线性关系的增强（最大到 1，其意味着完全线性关系）。负相关关系的降低意味着负（反向）线性关系的增强（最小到 -1，其意味着完全反向线性关系）。[⊖]

例 4-12 投资组合期望收益率及其方差

你有一个包含两个共同基金 A 和 B 的投资组合，其中 75% 投资于 A，如表 4-10 中所示。

(1) 计算投资组合的期望收益率。

(2) 计算本问题的相关系数矩阵，将结果保留到小数点后两位。

(3) 计算投资组合收益率的标准差。

解：

(1) $E(R_p) = w_A E(R_A) + (1 - w_A) E(R_B) =$

$0.75 (20\%) + (1 - 0.75) (12\%) = 18\%$ 。投资组合的权重必须相加为 1： $w_B = 1 - w_A$ 。

表 4-10 共同基金期望收益率、收益率方差和协方差

基金	A	B
	$E(R_A) = 20\%$	$E(R_B) = 12\%$
协方差矩阵		
基金	A	B
A	625	120
B	120	196

⊖ 如果相关系数为 0， $R_1 = a + bR_2 + \text{误差项}$ 中的 $b = 0$ 。

⊖ 如果相关系数为正， $R_1 = a + bR_2 + \text{误差项}$ 中的 $b > 0$ 。如果相关系数为负， $b < 0$ 。

(2) $\sigma(R_A) = 625^{1/2} = 25\%$, $\sigma(R_B) = 196^{1/2} = 14\%$ 。由于有一个不同的协方差, 因此, 有一个相关系数:

$$\rho(R_A, R_B) = \text{Cov}(R_A, R_B) / \sigma(R_A) \sigma(R_B) = 120 / 25 (14) = 0.342857 \text{ 或者 } 0.34。$$

表 4-11 列出了相关系数矩阵。

相关系数矩阵中对角项总是等于 1。

(3)
$$\begin{aligned} \sigma^2(R_p) &= w_A^2 \sigma^2(R_A) + w_B^2 \sigma^2(R_B) + 2w_A w_B \text{Cov}(R_A, R_B) \\ &= (0.75)^2 (625) + (0.25)^2 (196) \\ &\quad + 2(0.75)(0.25)(120) \\ &= 351.5625 + 12.25 + 45 = 408.8125 \\ \sigma(R_p) &= 40.8125^{1/2} = 20.22\% \end{aligned}$$

表 4-11 相关系数矩阵		
	A	B
A	1.00	0.34
B	0.34	1.00

我们如何估计收益率的协方差和相关系数? 通常, 我们基于历史协方差或者使用基于历史数据的其他方法 (如市场模型回归[⊖]) 来作出预测。我们也能利用随机变量的联合概率函数 (joint probability function) 来计算协方差 (如果它们都可以被估计出来)。两个随机变量 X 和 Y 的联合概率函数, 记为 $P(X, Y)$, 给出了 X 和 Y 的值同时发生的概率值。例如 $P(3, 2)$ 为出现 X 等于 3、 Y 等于 2 的概率值。

假设 BankCorp 股票收益率 (R_A) 和 NewBank 股票收益率 (R_B) 的联合概率函数具有表 4-12 中给出的简单结构。

BankCorp 股票的期望收益率为 0.20 (25%) + 0.50 (12%) + 0.30 (10%) = 14%。

NewBank 股票的期望收益率为 0.20 (20%) + 0.50 (16%) + 0.30 (10%) = 15%。上面的联合概率函数可能反映了一个基于银行业情况是良好、中等还是不佳的分析。表 4-13 给出了协方差的计算。

表 4-12 BankCorp 和 NewBank 收益率的联合概率函数 (表中元素为联合概率)

	$R_B = 20\%$	$R_B = 16\%$	$R_B = 10\%$
$R_A = 25\%$	0.20	0	0
$R_A = 12\%$	0	0.50	0
$R_A = 10\%$	0	0	0.30

表 4-13 协方差的计算

银行业的情况	BankCorp 的偏差值	NewBank 的偏差值	偏差值的乘积	出现该情况的概率	概率加权乘积
良好	25 ~ 14	20 ~ 15	55	0.20	11
中等	12 ~ 14	16 ~ 15	-2	0.50	-1
不佳	10 ~ 14	10 ~ 15	20	0.30	6
					$\text{Cov}(R_A, R_B) = 16$

注: BankCorp 的期望收益率为 14%, 而 NewBank 的期望收益率为 15%。

第 1 栏和第 2 栏中的数分别显示出 BankCorp 和 NewBank 收益率偏离它们均值或者期望值的偏差程度。下面一栏显示的是偏差值的乘积。例如, 对于良好的行业情况, $(25 - 14)(20 - 15) = 11(5) = 55$ 。接着 55 乘以 0.20 或者以 0.20 (银行业的情况为良好的概率) 作为权重, 结果为 $55(0.20) = 11$ 。银行业的情况是中等或不佳时的计算方法与之相同。将这些概率加权乘积进行加总, 我们得到 $\text{Cov}(R_A, R_B) = 16$ 。

计算随机变量 R_A 和 R_B 之间协方差的计算公式为:

⊖ 参见脚注 10 中所列举的教科书。

$$\text{Cov}(R_A, R_B) = \sum_i \sum_j P(R_{A,i}, R_{B,j}) (R_{A,i} - ER_A)(R_{B,j} - ER_B) \quad (4-18)$$

公式告诉我们将所有可能偏差值以适当联合概率加权交叉乘积求和。在上面所做的这个例子中，如表 4-12 所示，只有 3 个联合概率是非零的。因此，在这个例子的收益率协方差计算过程中，我们只需要考虑 3 个交叉乘积：

$$\begin{aligned} \text{Cov}(R_A, R_B) &= P(25, 20) [(25 - 14)(20 - 15)] + P(12, 16) [(12 - 14)(16 - 15)] \\ &\quad + P(10, 10) [(10 - 14)(10 - 15)] \\ &= 0.20(11)(5) + 0.50(-2)(1) + 0.30(-4)(-5) \\ &= 11 - 1 + 6 = 16 \end{aligned}$$

本章的一个主题是独立性。当每一可能的事件对（对应于 X 值的一个事件和对应于 Y 值的另一个事件）是独立事件时，两个随机变量是独立的。当两个随机变量是独立的，它们的联合概率函数可以得到简化。

- **随机变量独立性的定义。**两个随机变量 X 和 Y 是独立的当且仅当 $P(X, Y) = P(X)P(Y)$ 。

例如，给定两个随机变量独立， $P(3, 2) = P(3)P(2)$ 。我们将单个概率相乘得到联合概率。独立性是比不相关性更强的性质，因为相关性只是表现了线性关系。下面的条件对于独立的随机变量成立，因此，对于不相关的随机变量也同样成立。

- **不相关随机变量乘积期望值的乘法法则**（multiplication rule for expected value of the product of uncorrelated random variables）。不相关随机变量乘积的期望值为它们期望值的乘积。

当 X 和 Y 是不相关时， $E(XY) = E(X)E(Y)$ 。

许多金融变量，如收入（价格乘以数量），是随机量的乘积。在应用时，上面的法则简化了计算随机变量乘积期望值的计算。[Ⓐ]

4.4 概率论的一些议题

在本章的剩余部分，我们将讨论两个可能对于求解投资问题非常重要的主题。我们从贝叶斯公式开始，该概率理论是关于如何从经验中进行学习。接着，我们将转向计算中诀窍以及计数原理的讨论。

4.4.1 贝叶斯公式

当我们做出一个关于投资的决策时，我们经常会基于我们的经验和知识形成最初的观点。这些观点可能根据新的信息或者观察而发生改变或者加以确认。贝叶斯公式是在我们遇到新信息时调整我们观点的一种理性方法。[Ⓑ] 贝叶斯公式和相关的概念已经被应用于许多商业和投资的决策环境之中，包括对于共同基金的业绩评价。[Ⓒ]

贝叶斯公式利用式（4-6）的全概率公式。回顾一下，该法则将一个事件的概率表达为给定一组其情景条件下的事件概率的加权平均值。贝叶斯公式以相反的方式运作；更确切地说，它逆

Ⓐ 不然，计算依赖于条件期望值；计算将被表达成 $E(XY) = E(X)E(Y | X)$ 。

Ⓑ 以教士托马斯·贝叶斯（Thomas Bayes（1702—1761））命名。

Ⓒ 参见 Baks, Metrick 和 Wachter（2001）。

转了“给定的”信息。贝叶斯公式在事件发生的情况下来推断产生它的这一情景的概率。因为这个原因，贝叶斯公式有时也被称为逆概率。在许多应用中，包括在这一节中所介绍的例子，人们会更新其关于引起新观测值原因的信念。

- 贝叶斯公式 (Bayes' formula)。给定一组所关心事件的先验概率，如果你收到新的信息，那么更新你对于事件发生概率的法则为：

在给定新信息下更新的事件发生的概率 = $\frac{\text{给定事件新信息发生的概率}}{\text{新信息发生的无条件概率}} \times \text{事件发生的先验概率}$

用概率的记号，该公式可以简单地写成：

$$P(\text{事件} \mid \text{信息}) = \frac{P(\text{信息} \mid \text{事件})}{P(\text{信息})} P(\text{事件})$$

为了举例说明贝叶斯公式，我们要做一个适用于任何实际问题的投资的例子。假设你是 DriveMed 公司股票的投资者。意料之外的相对于人们共同的 EPS 估计值的正的盈利常常导致正的股票收益率，而意料之外的负面情况经常产生相反的结果。DriveMed 只能公布最近一季度的 EPS 结果，而你关注于这 3 件事件哪个会发生：最近一季度的 EPS 超出人们共同的 EPS 估计值，或者最近一季度的 EPS 与人们共同的 EPS 估计值相等，又或者最近一季度的 EPS 估计值少于人们共同的 EPS 估计值。这列备选事件是互斥且穷举的。

基于你的研究，你写下了如下关于这 3 个事件的先验概率 (prior probabilities)，或者简要地称为先验值 (priors)：

$$\begin{aligned} P(\text{EPS 超过人们共同预期}) &= 0.45 \\ P(\text{EPS 等于人们共同预期}) &= 0.30 \\ P(\text{EPS 小于人们共同预期}) &= 0.25 \end{aligned}$$

这些概率是先验的，因为它们反映了在新信息到来之前，你现在所知道的信息。

第 2 天，DriveMed 公布其为了满足增长的销售需求在新加坡和爱尔兰扩大其工厂规模的消息，你获得了这一新的信息。扩大规模的决定不仅与当前的需求有关，而且可能与先验的季度销售需求有关。你知道销售需求是与 EPS 正相关的。所以现在来看，下个季度的 EPS 似乎更可能超出人们共同的预期。

你现在的的问题是，“在新信息出现的情况下，先前的下个季度的 EPS 更可能会超出人们共同预期的更新概率是多少？”

贝叶斯公式提供了完成这一更新的一个理性方法。我们简要地把信息记为 DriveMed 扩张。应用贝叶斯公式的第一步是在给定一系列产生它的事件或者情景的情况下，计算新信息出现的概率 (这里指 DriveMed 扩张的概率)。该列事件应该包括所有概率，正如下面所列出的。写出这些条件概率是更新过程中的重要一步。假设你的观点是：

$$\begin{aligned} P(\text{DriveMed 扩张} \mid \text{EPS 超过人们共同预期}) &= 0.75 \\ P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) &= 0.20 \\ P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) &= 0.05 \end{aligned}$$

一个观测 (这里指 DriveMed 扩张) 的条件概率有时被称为似然值 (可能性) (likelihoods)。此外，似然值是需要用来更新概率的。

下面，你要将这些条件概率或者似然值与你的先验概率整合起来以获得 DriveMed 扩张的无条

件概率, $P(\text{DriveMed 扩张})$ 的计算如下:

$$\begin{aligned} P(\text{DriveMed 扩张}) &= P(\text{DriveMed 扩张} \mid \text{EPS 超过人们共同预期}) \\ &\quad \times P(\text{EPS 超过人们共同预期}) + P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) \\ &\quad \times P(\text{EPS 等于人们共同预期}) + P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) \\ &\quad \times P(\text{EPS 小于人们共同预期}) = 0.75(0.45) \\ &\quad + 0.20(0.30) + 0.05(0.25) = 0.41 \text{ 或者 } 41\% \end{aligned}$$

这是式 (4-6) 的全概率公式在起作用。现在你可以用贝叶斯公式回答你的问题了:

$$\begin{aligned} &P(\text{EPS 等于人们共同预期} \mid \text{DriveMed 扩张}) \\ &= \frac{P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期})P(\text{EPS 等于人们共同预期})}{P(\text{DriveMed 扩张})} \\ &= (0.75/0.41)(0.45) = 1.829\,268(0.45) = 0.823\,171 \end{aligned}$$

在 DriveMed 公布公告之前, 你认为 DriveMed 会超出人们共同预期的可能性为 45%。而在基于你对于公告的解释之后, 你将该概率值更新为 82.3%。这个更新概率被称为后验概率 (posterior probability), 因为它反映了新信息或者在新信息到来之后所得到的概率。

贝叶斯的计算方法利用了先验概率, 这里是 45%, 再将其乘以一个比率 (等号右边的第 1 项)。该比率分母是 DriveMed 扩张的概率, 正如你所看到的, 它不考虑任何条件 (不基于任何条件)。因此, 该概率是无条件概率。其分子是如果最近一个季度 EPS 实际超出人们共同预期情况下 DriveMed 扩张的概率。这个概率要大于分母无条件概率, 故比率 (约为 1.83) 比 1 要大。结果, 你更新的或者后验的概率比你的先验概率要大。因此, 比率反映了新信息对于你先前信念的影响。

例 4-13 推断 DriveMed 的 EPS 是否会符合人们共同的预期

你仍旧是持有 DriveMed 股票的一个投资者。回顾一下之前的已知条件, 你的先验概率为 $P(\text{EPS 超出人们共同预期}) = 0.45$, $P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) = 0.30$, $P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) = 0.25$ 。你也获得了下列条件概率:

$$P(\text{DriveMed 扩张} \mid \text{EPS 超过人们共同预期}) = 0.75$$

$$P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) = 0.20$$

$$P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) = 0.05$$

回忆一下, 你在 DriveMed 公布其将扩张的消息之后, 将关于最近一季度 EPS 超出人们共同预期的概率从 45% 更新到了 82.3%。现在你想要更新你的其他先验值。

(1) 更新你关于 DriveMed 的 EPS 等于人们共同预期的先验概率。

(2) 更新你关于 DriveMed 的 EPS 小于人们共同预期的先验概率。

(3) 表明这 3 个更新的概率之和等于 1。(每个概率值保留到小数点后 4 位。)

(4) 假如因为缺乏关于 DriveMed 是否会符合人们共同的预期的先验信念, 你在对所有 3 种可能性以等概率发生的基础上进行概率的更新: $P(\text{EPS 超出人们共同预期}) = P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) = P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) = 1/3$ 。那么你对于概率 $P(\text{EPS 超过人们共同预期} \mid \text{DriveMed 扩张})$ 的估计值又是多少?

解:

(1) 概率为

$$\begin{aligned} & P(\text{EPS 等于人们共同预期} \mid \text{DriveMed 扩张}) \\ &= \frac{P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期})}{P(\text{DriveMed 扩张})} P(\text{EPS 等于人们共同预期}) \end{aligned}$$

概率 $P(\text{DriveMed 扩张})$ 通过将问题中所论述的 3 个条件概率中的每个 (如 $P(\text{DriveMed 扩张} \mid \text{EPS 超出人们共同预期})$) 乘以每个条件事件的先验概率 (如 $P(\text{EPS 超出人们共同预期})$), 再将这 3 个乘积相加得到。在上面的问题中, 该计算方法不变: $P(\text{DriveMed 扩张}) = 0.75(0.45) + 0.20(0.30) + 0.05(0.25) = 0.41$ 或者 41%。其他所需的概率, $P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) = 0.20$ 和 $P(\text{EPS 等于人们共同预期}) = 0.30$ 是已知的。所以

$$\begin{aligned} P(\text{EPS 等于人们共同预期} \mid \text{DriveMed 扩张}) &= [P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) / \\ &P(\text{DriveMed 扩张})] P(\text{EPS 等于人们共同预期}) = (0.20/0.41)(0.30) = 0.487805(0.30) = \\ &0.146341 \end{aligned}$$

在考虑扩张公告之后, 与你之前先验概率 30% 相比, 你对于最近一季度 DriveMed 的 EPS 正好符合人们共同预期的更新概率为 14.6%。

(2) $P(\text{DriveMed 扩张})$ 已经计算出, 其值为 41%。回忆一下, $P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) = 0.05$ 以及 $P(\text{EPS 小于人们共同预期}) = 0.25$ 是已知的。

$$\begin{aligned} & P(\text{EPS 小于人们共同预期} \mid \text{DriveMed 扩张}) \\ &= [P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) / P(\text{DriveMed 扩张})] \\ &P(\text{EPS 小于人们共同预期}) = (0.05/0.41)(0.25) \\ &= 0.121951(0.25) = 0.030488 \end{aligned}$$

结果在公告发布之后, 你将 DriveMed 的 EPS 小于人们共同预期的概率从 25% (你的先验概率) 修改为 3%。

(3) 这 3 个更新概率的总和为

$$\begin{aligned} & P(\text{EPS 超出人们共同预期} \mid \text{DriveMed 扩张}) \\ &+ P(\text{EPS 等于人们共同预期} \mid \text{DriveMed 扩张}) \\ &+ P(\text{EPS 小于人们共同预期} \mid \text{DriveMed 扩张}) \\ &= 0.8232 + 0.1463 + 0.0305 = 1.0000 \end{aligned}$$

这 3 个事件 (EPS 超出人们共同预期, EPS 等于人们共同预期, EPS 小于人们共同预期) 是互斥且穷举的: 这些事件或者叙述之一必须出现, 所以条件概率必须加总为 1。无论我们讨论条件概率还是无条件概率, 当我们有完整的互不相同的一组可能事件或者结果, 其概率之和必须为 1。这一计算可以作为对于计算结果的一个检验。

(4) 利用问题中给定的概率

$$\begin{aligned} P(\text{DriveMed 扩张}) &= P(\text{DriveMed 扩张} \mid \text{EPS 超出人们共同预期}) \\ &P(\text{EPS 超出人们共同预期}) + P(\text{DriveMed 扩张} \mid \text{EPS 等于人们共同预期}) \\ &P(\text{EPS 等于人们共同预期}) + P(\text{DriveMed 扩张} \mid \text{EPS 小于人们共同预期}) \\ &P(\text{EPS 小于人们共同预期}) = 0.75(1/3) + 0.20(1/3) + 0.05(1/3) \\ &= 1/3 \end{aligned}$$

毫不奇怪, DriveMed 扩张的概率为 $1/3$, 因为决策者没有先验的关于 EPS 相对于共同预期将如何表现的信念或者观点。现在我们可以利用贝叶斯公式来求解 $P(\text{EPS 超过人们共同预期} | \text{DriveMed 扩张}) = [P(\text{DriveMed 扩张} | \text{EPS 超过人们共同预期}) / P(\text{DriveMed 扩张})] P(\text{EPS 超过人们共同预期}) = [(0.75) / (1/3)] (1/3) = 0.75$ 或者 75%。这一概率等于你对于 $P(\text{DriveMed 扩张} | \text{EPS 超过人们共同预期})$ 的估计。

当先验概率相等的时候, 给定事件信息发生的概率等于给定信息事件发生的概率。当一个决策者具有相等的概率 (称为原始先验概率假设 (diffuse priors)), 一个事件的概率是由信息所决定的。

4.4.2 计数原理

叙述一个问题的第一步经常涉及确定不同情况的概率, 我们可能也想知道产生每种这样结果的方式有多少。在我们的思维中, 总是存在一个关于概率的问题。我们观测到这一特定概率的可能性到底是多少? 对于成功和失败的记录就是一个例子。当我们评估一个市场择时者的记录时, 所采用的一个著名方法就是本节中所介绍的计数方法。^① 一个重要的投资模型——二叉树期权定价模型包含了我们不久将讨论的组合公式。我们也可以使用本节中的方法来计算在第2节中的先验概率。当我们认为随机变量的可能结果相等是可能的, 那么一个事件发生的概率等于对于该事件有利的可能结果的数目除以结果的总数。

在计数时, 枚举法 (逐个计算每个结果) 当然是最基本的方法, 但我们在本节中所讨论的是一些诀窍和原理。除了这些诀窍和原理, 计算结果的总数可能非常困难而且易于出错。第一个也是最基本的计数原理是乘法法则。

- 计数的乘法法则 (multiplication rule of counting)。如果完成第1项工作有 n_1 种方法; 在第1项工作完成之后, 完成第2项工作有 n_2 种方法; 在完成这两项工作之后, 完成第3项工作有 n_3 种方法, 以此类推, 直到第 k 项工作, 于是完成这 k 项工作共有 $(n_1)(n_2)(n_3) \cdots (n_k)$ 。

假设我们的一个投资决策过程有3个步骤。完成第1步有两种方法, 第2步有4种方法, 第3步有3种方法。那么根据乘法法则, 我们完成这3个步骤共有 $(2)(4)(3) = 24$ 种方法。

另外一个问题是将一组元素分配到相等数目的位置之中。例如, 假设你想要分配3名证券分析师去研究3个不同的行业, 那么共有多少种分配方法? 第1个分析师可能有3种不同的分配方式。接着剩下两个行业。故第2个分析师有两种分配方式。于是只剩一个行业。第3个也是最后一个分析师只有一种选择。故不同分配的种数为 $(3)(2)(1) = 6$ 。将该乘积以简单记号来表达就是去求解 $3!$ (读作3的阶乘)。如果我们有 n 位分析师, 那么我们将他们分配到 n 个工作岗位上的种数为

$$n! = n(n-1)(n-2)(n-3) \cdots 1$$

或者 n 的阶乘 (n factorial)。(通常我们有 $0! = 1$ 。) 回顾一下, 在这个应用中我们重复进行操作 (这里是工作的分配) 直到我们用尽该组中的所有元素 (这里是3位分析师)。当组里有 n 个元素

① Henriksson 和 Merton (1981)。

时,乘法公式将变为 n 的阶乘。[⊖]

下一类的计数问题可以被称为贴标签问题。[⊗] 我们想要给组内的每个元素贴上一个标签,即将它分配一个类别之中。下面的例子说明了这一类型的问题。

一位共同基金投资顾问将 18 只债券型共同基金按照 2000 年总收益率情况进行排序。该顾问同时将每个基金分别归到 5 个风险类别之中:高风险(4 个基金),高于平均风险(4 个基金),平均风险(3 个基金),低于平均风险(4 个基金),低风险(3 个基金)。所有基金共有 $4 + 4 + 3 + 4 + 3 = 18$ 家。那么我们将这 18 家基金按照高风险 4 家、高于平均风险 4 家、平均风险 3 家、低于平均风险 4 家、低风险 3 家的方式进行分类,以使得每个基金都有各自的类别,共有多少种不同的方法?

该问题的答案接近 130 亿。我们可以将 18 家基金中的任何一家归为高风险(第 1 顺位),接着剩下 17 家基金中的任何一家归为高风险,再接着剩下 16 家基金中的任何一家归为高风险,最后剩下 15 家基金中的任何一家归为高风险(现在我们将 4 家基金归到了高风险组中);接着我们可以将剩下 14 家基金中的任何一家归为高于平均风险类中,接着将剩下 13 家中的任何一家归为高于平均风险类中,以此类推。于是共有 $18!$ 种可能的序列。然而,分配到类别中的先后次序并没有影响。例如,无论一家基金是第一顺位分配到高风险类别中的 4 家基金之一还是第 3 顺位分配到高风险类别中,最后该基金总具有相同的类别(高风险类)。因此,将给定的一组 4 家基金分配到高风险类的 4 个顺位之中,共有 $4!$ 种方法。对于其他类别做出相同的论述,总共有 $(4!)(4!)(3!)(4!)(3!)$ 种等价序列。为了从 $18!$ 这个总数中消除重复的计数,我们将 $18!$ 除以 $(4!)(4!)(3!)(4!)(3!)$ 。我们有 $18! / (4!)(4!)(3!)(4!)(3!) = 18! / (24)(24)(6)(24)(6) = 12\,864\,852\,000$ 。该过程可以一般化为:

- 多项式公式(贴标签问题的一般公式)(multinomial formula (general formula for labeling problems))。给 n 个元素标上 k 种不同的标签,第一类有 n_1 个,第 2 类有 n_2 个,以此类推,共有 $n_1 + n_2 + \cdots + n_k = n$, 其分配的种数为

$$\frac{n!}{n_1!n_2!\cdots n_k!}$$

两个不同标签($k=2$)的多项式公式尤其重要。该特殊情形被称为组合公式。组合是一种选取方法,所选取元素的先后顺序没有关系。我们以一种传统的方式来表达组合公式,但是并没有引入新的概念。利用下面公式中的记号,标为第一个标签的元素有 $r=n_1$ 个,标为第 2 个标签的元素有 $r=n_2$ 个(由于只有两个类别,故 $n_1 + n_2 = n$)。下面是组合公式。

- 组合公式(二项式公式)(combination formula (binomial formula))。当列出的 r 个元素的先后顺序没有差别时,我们从 n 个元素中选出 r 个元素的种数为

$${}_nC_r = \binom{n}{r} = \frac{n!}{(n-r)!r!}$$

这里 ${}_nC_r$ 和 $\binom{n}{r}$ 都是 $\frac{n!}{(n-r)!r!}$ 的缩略记号(读做 n 中选 r 个或者 n 选 r 的组合数)。

⊖ n 阶乘的简要解释为将 n 个元素排在一行的种数。在所有的应用这种计数方法的问题中,我们必须用尽一组中所有的元素(无重复抽样)。

⊗ 该讨论沿袭了 Kemeny, Schleifer, Snell 和 Thompson (1972) 所使用的术语与方法。

如果我们将 r 个元素标为属于该类别而剩余其他元素不属于该类别,那么无论我们所关注的类别是什么,组合公式告诉我们选择大小为 r 的类别的种数共有多少。我们可以用二叉树期权定价模型来说明这个公式。该模型将标的资产价格的运动描述为一系列上升 (U) 或者下降 (D) 的价格移动。例如,5 次移动中包含 3 次上升的两个序列为 UUDD 和 UDUUD,它们的最终股票价格相同。至少对于依赖于期末股票价格收益的期权价格,不依赖序列中股价上升的排列先后顺序。5 次移动中包含 3 次上升的序列共有多少种? 回答是 10,利用组合公式计算如下 (“5 选 3”):

$${}_5C_3 = \frac{5!}{(5-3)!3!} = \frac{(5)(4)(3)(2)(1)}{(2)(1)(3)(2)(1)} = \frac{120}{12} = 10 \text{ 种}$$

一个有用的事实可以阐述如下: ${}_5C_3 = 5! / 2! 3!$ 等于 ${}_5C_2 = 5! / 3! 2!$, 等于 $3 + 2 = 5$; ${}_5C_4 = 5! / 1! 4!$ 等于 ${}_5C_1 = 5! / 4! 1!$, 等于 $4 + 1 = 5$ 。当我们需要计算许多可能的组合数时,利用这一对称关系可以节省计算时间。

假设评审团想要从一组 5 家公司中选出 3 家公司以评出最佳年报的一、二、三等奖,那么评审团可以有多少种评奖的结果? 如果我们要区分这 3 个奖项间的不同 (这 3 个奖项的排名情况),那么排列的先后次序确实和结果有关;显然该问题使得排序非常重要。而另一方面,如果该问题是“评审团选 3 个胜者而不考虑最终的先后顺序?”我们会用组合公式来进行求解。

为了回答上面的第一个问题,我们需要计算诸如第 1 顺位 New 公司、第 2 顺位 Fir 公司、第 3 顺位 Well 公司这样的顺序列表。一个排序列表是一个排列 (permutation),计算排列数目的公式是排列公式。[⊖]

- 排列公式 (permutation formula)。我们从总数为 n 个元素中选取 r 个元素,而且要考虑 r 个元素的排列顺序,其总共的排列方法为

$${}_nP_r = \frac{n!}{(n-r)!}$$

所以评审团有 ${}_5P_3 = 5! / (5-3)! = \frac{(5)(4)(3)(2)(1)}{(2)(1)} = \frac{120}{2} = 60$ 种排列它们奖项的方法。

为了观察这个公式为什么是正确的,注意到 $(5)(4)(3)(2)(1)$ 在约去共同项之后,为 $(5)(4)(3)$ 。该计算根据计数的乘法法则来计算从 5 个人的组中选出 3 个的方法总数。该数字自然要比顺序先后不产生影响所得到的结果要大 (将 60 与我们上面 “5 选 3” 计算出的 10 种方法相比)。例如,第 1 位为 Well 公司;第 2 位为 Fir 公司;第 3 位为 New 公司的结果与 3 个公司排列为第 1 位为 New 公司、第 2 位为 Fir 公司、第 3 位为 Well 公司的结果相同。如果我们只考虑得奖的赢家 (不考虑最终结果的顺序),那么这两种排列将视为一种组合。但是当我们考虑最终结果的顺序时,这两种排列就是两个排列。

回答下列问题可以帮助你应用本节中所介绍的计数方法。

- (1) 我想要度量的问题是否有有限个可能结果? 如果回答是肯定的,你可以利用本章中的工具,并继续回答第 2 个问题。如果回答是否定的,即结果的数目是无穷的,则不能使用本节中的工具。

⊖ 更正式的定义为,一个排列是 n 个不同元素排序的子集。

- (2) 我是否要将大小为 n 组中的每个元素分配到 n 个位置（或者任务）之中？如果回答是可能的，则使用 n 的阶乘。如果回答是否定的，则转到第3个问题。
- (3) 我是否要计算将3个或者更多的标签分配给组内的每个元素？如果回答是肯定的话，利用多项式公式。如果回答是否定的，则转到第4个问题。
- (4) 我是否要计算从 n 个总体中选取 r 个元素的种数，并且我们选出的 r 个元素的先后次序没有影响（我能否给 r 个元素中的每个一个不同标签）？如果这些问题的回答是肯定的，那么应用组合公式。如果回答是否定的，转到第5个问题。
- (5) 我是否要计算从 n 个总体中选出 r 个元素的种数，而且选出的 r 个元素的先后次序非常重要？如果回答是肯定的，那么应用排列公式。如果回答是否定的，转到第6个问题。
- (6) 是否能够应用计数的乘法法则？如果不能，你可能需要逐个计算出问题的概率，或者使用比这里介绍的更加高级的技术。[⊖]

⊖ Feller (1957) 包含了非常完整的计数问题的处理和解决方法。



常用概率分布

5.1 引言

几乎在所有的投资决策中，我们都会用到随机变量。股票收益率和每股收益都是我们非常熟悉的随机变量的例子。为了能给出一个随机变量发生概率情况的表述，我们需要了解随机变量的概率分布。一个概率分布（probability distribution）具体规定了一个随机变量可能出现的所有结果的概率。

在这一章中，我们将给出4个概率分布的重要性质以及它们在投资中的应用。这4个概率分布（均匀分布、二项分布、正态分布和对数正态分布）被广泛应用于投资分析之中。像布莱克－斯科尔斯－默顿期权定价模型、二叉树期权定价模型以及资本资产定价模型中都有这些分布的应用。有了本章中所提供的概率分布的实用知识，你能很好地奠定进一步学习和应用其他定量模型如假设检验、回归分析和时间序列分析的基础。

在讨论完概率分布之后，我们在本章的结尾处将介绍一下蒙特卡罗模拟，这是一种利用计算机来获得复杂问题信息的方法。例如，一位投资分析师可能想要测试一下一个投资想法而不是真正地在现实中实施。或者她可能需要为一个复杂的期权定价，而该期权不存在简单的定价公式。在这些例子和许多其他的例子中，蒙特卡罗模拟是一个重要的办法。为了进行一个蒙特卡罗模拟，分析师必须识别与问题相关的风险因素，并确定它们的概率分布。因此，蒙特卡罗模拟是一项需要了解概率分布的工具。

在我们开始讨论具体的概率分布之前，我们先定义一些基本的概念和常用的术语。然后，我们再通过最简单的分布（均匀分布）来介绍这些概念的应用。在完成这些内容之后，我们再讨论一些在投资工作中有更多应用而且更加复杂的概率分布。

5.2 离散型随机变量

随机变量（random variable）是一个未来发生的结果不确定的量。两种基本的随机变量类别

为离散型随机变量和连续型随机变量。一个离散型随机变量 (discrete random variable) 至多可以取可数个可能的结果值。例如, 离散型随机变量 X 可以取有限个结果值 $x_1, x_2, \dots, x_n$ (n 个可能的结果), 或者一个离散型随机变量 Y 可能取无限多个结果值 $y_1, y_2, \dots$ (无终结的)。^① 因为我们可以数出 X 和 Y 的所有可能的结果 (即使我们会在 Y 的例子中趋于无穷), 所以 X 和 Y 满足一个离散型随机变量的定义。相反, 我们不能够数出一个连续型随机变量 (Continuous random variable) 的所有结果。我们不能用一系列 $z_1, z_2, \dots$ 来描述一个连续型的随机变量 Z , 因为可能的结果 $(z_1 + z_2)/2$ 不在那个序列之中。收益率就是一个连续型随机变量的例子。

在处理一个随机变量时, 我们需要了解其可能发生的结果。例如, 在纽约证券交易所和纳斯达克上交易的股票是以 0.01 美元的最小变化进行报价的。因此, 报出的股票价格是具有 0 美元, 0.01 美元, 0.02 美元, $\dots$ 可能结果的离散型随机变量。但是我们也可以将股票价格以连续型随机变量进行建模 (作为一个对数正态随机变量来对未来进行预测)。在许多应用中, 我们需要在使用离散分布或者使用连续分布之间进行选择。我们经常把哪种分布对于我们的工作最方便作为选择的标准。这样选择的标准并不让人惊讶, 因为许多离散分布可以看做连续分布的近似, 反之亦然。在大部分实际的例子中, 一个概率分布只是一个对于随机变量可能发生结果的相对频率在数学上理想化的东西, 或者说只是一个近似的模型。

例 5-1 债券价格的分布

你正在研究一个债券价格的概率模型。你从影响债券价格的特征开始考虑。债券价格最小和最大的可能值为多少? 为什么? 其他影响债券价格分布的债券特征是什么?

债券价格的最低可能值是 0, 这时的债券是不值钱的。确定债券价格最高可能值更困难一些。付息债券的承诺支付是息票 (利息支出) 加上面值 (本金)。债券的价格是这些承诺支付的贴现值。因为投资者对他们的投资要求一个回报率, 0% 是投资者用来折现债券承诺支付折现率的下极限。在折现率为 0% 的情况下, 债券价格是面值和所有没有经过贴现的剩余息票之和。因此, 折现率给出了债券价格的上极限。例如, 假设 (一个债券的) 面值为 1 000 美元, 而剩余两个息票分别为 40 美元; 区间 0 ~ 1 080 美元包括了该债券价格的所有可能值。这个上极限会随着时间的流逝而降低, 因为剩余支付的数量减少了。

其他债券特征也会影响债券价格的分布。回归面值具有这样的一种特征: 随着到期日的接近, 随着债券价格收敛到面值, 债券价格的标准差倾向于变小。嵌入式期权也会影响债券价格。例如对于当前可赎回的债券, 发行者可以以一个事先设定的高于面值的价格将债券赎回; 这一发行者的选择权消除了债券价格上行部分结果的出现。对于这类债券价格分布建模是一项具有挑战性的课题。

每个随机变量是与一个能完全描述该变量的概率分布相联系的。我们能从两个方面来看待这个概率分布。一个基本的观点是把它看做一个概率函数 (probability function)。该函数定义了随机变量取一个特定值的概率: $P(X = x)$ 是随机变量 X 取 x 这个值时的概率。(注意大写的 X 表示随机变量, 而小写的 x 表示随机变量所取的一个特定的值。) 对于一个离散型随机变量来说, 其概率函数可以简写为 $p(x) = P(X = x)$ 。而对于一个连续型随机变量来说, 概率函数

① 我们采用传统的记法: 大写字母表示一个随机变量, 小写字母表示该随机变量出现的一个结果或者具体的一个数值。因此, X 代表随机变量, 而 x 代表 X 的一个结果。当我们需要区分一个随机变量的不同结果时, 我们给每个结果标上下标, 如 x_1 和 x_2 。

记为 $f(x)$ ，并称其为概率密度函数（probability density function, pdf），或者简称为概率密度。[⊖]
一个概率分布函数具有两个重要的性质（不失一般性，我们用一个离散型随机变量的记号来进行说明）：

- $0 \leq p(x) \leq 1$ ，因为概率是一个介于0~1之间的数值。
- X 所有可能取值的概率 $p(x)$ 之和等于1。如果我们将一个随机变量所有不同可能结果的概率值相加，其总和必须等于1。

我们经常关注求解出现在一定范围的结果的概率值，而不是一个特定结果的概率值。在这种情况下，我们要从第二个角度来看待概率分布，即把它理解为一个累积分布函数（cumulative distribution function, cdf）。累积分布函数或者简称分布函数，给出了一个随机变量小于等于某个特定值 x 的概率 $P(X \leq x)$ 。对于离散型和连续型随机变量，我们都把分布函数记为 $F(x) = P(X \leq x)$ 。那么累积分布函数是怎样和概率函数联系起来的呢？“累积”一词给出了答案。为了求解 $F(x)$ ，我们需要加总或者累积所有小于等于 x 这一结果的概率函数值。cdf函数和累积相对频率（这个概念我们在统计学概念和市场收益率一章中已经进行了讨论）是相对应的。

下面，我们将用例子来举例说明和展现如何使用这些离散和连续的分布。我们首先从最简单的分布——离散均匀分布开始。

5.2.1 离散均匀分布

所有概率分布中最简单的概率分布就是离散均匀分布。假设可能出现的结果是整数（integer, whole number），从1~8（包含端点的），那么该随机变量取任何可能值的概率对于每个结果来说是相同的（即它是均匀的）。由于有8个结果，故 $p(x) = 1/8$ 或者0.125（对于所有的 $X = 1, 2, 3, 4, 5, 6, 7, 8$ ）；刚才的表述是对于这个离散型均匀随机变量的一个完整描述。分布具有有限个特定结果，且每个结果是等可能性的。表5-1从概率函数和累积分布函数两种角度对该随机变量进行了描述。

我们利用表5-1来计算3个概率： $P(X \leq 7)$ ， $P(4 \leq X \leq 6)$ 和 $P(4 < X \leq 6)$ 。下面的例子说明了如何使用cdf来计算一个随机变量落在某一区间中的概率（对于任何随机变量，而不仅仅是均匀随机变量）。

X 小于等于7的概率 $P(X \leq 7)$ 等于在第3栏倒数第2行格子中的数值0.875或者87.5%。

为了计算 $P(4 \leq X \leq 6)$ ，我们需要计算3个概率的总和： $p(4)$ ， $p(5)$ 和 $p(6)$ 。我们可以用两种方法来求解该总和。我们可以加总第2栏中的 $p(4)$ 、 $p(5)$ 和 $p(6)$ ，或者我们可以通过两个累积分布函数值之差来计算其概率：

$$F(6) = P(X \leq 6) = p(6) + p(5) + p(4) + p(3) + p(2) + p(1)$$
$$F(3) = P(X \leq 3) = p(3) + p(2) + p(1)$$

于是 $P(4 \leq X \leq 6) = F(6) - F(3) = p(6) + p(5) + p(4) = 3/8$

表 5-1 对于一个离散型均匀随机变量的
概率分布函数和累积分布函数

$X = x$	概率分布函数 $p(x) = P(X = x)$	累积分布函数 $F(x) = P(X \leq x)$
1	0.125	0.125
2	0.125	0.250
3	0.125	0.375
4	0.125	0.500
5	0.125	0.625
6	0.125	0.750
7	0.125	0.875
8	0.125	1.000

⊖ 一个离散随机变量概率函数的专业术语是概率质量函数（pmf），但是这种说法很少使用。

因此，我们计算的第2个概率为 $F(6) - F(3) = 3/8$ 。

第3个概率， $P(4 < X \leq 6)$ ， X 小于等于6但是大于4的概率为 $p(5) + p(6)$ 。我们利用 cdf 将其计算如下：

$$P(4 < X \leq 6) = P(X \leq 6) - P(X \leq 4) = F(6) - F(4) = p(5) + p(6) = 2/8$$

所以我们计算出第3个概率为 $F(6) - F(4) = 2/8$ 。

假设我们想要检验该离散均匀概率函数是否满足之前给出的概率分布函数的一般性质。第一个性质 $0 \leq p(x) \leq 1$ 。我们可以看到 $p(x) = 1/8$ 对于所有表中第一栏中的 x 成立。（注意对于 x 等于 -14 或者 12.215 这类值的 $p(x)$ 等于0，因为这些值不在第一栏中。）故第一个性质是满足的。第2个性质是概率之和等于1。表5-1中第2栏中所有的元素之和确实相加为1。

cdf 具有另外两个性质：

- 对于任何 x ，cdf 位于0和1之间： $0 \leq F(x) \leq 1$ 。
- 随着 x 的增加，cdf 或者增加或者不变。

我们可以通过观察表5-1中第3列中的数据来检验这两个表述。

我们现在已经有了利用概率函数和 cdf 来处理离散型随机变量的一些经验。在本章的后面，我们将讨论蒙特卡罗模拟——一种基于随机数量的处理方法。正如我们将会看到的，均匀分布具有一个重要的技术上的应用：它是产生随机数的基础，而随机数将会最终产生所有其他概率分布的随机观察值。[⊖]

5.2.2 二项分布

在许多投资环境中，我们将可能出现的结果定义为成功或者失败，或者我们也可能在其他的一些情况下用二进制数（或2个不同的数值）来对其进行表示。当我们对于成功与失败的记录或者对于二元的结果做出概率表述时，我们经常使用二项分布。对于刻画股价如何随时间变动的最好模型是什么？不同的模型适合于不同的使用情况。考克斯（Cox）、罗斯（Ross）和鲁宾斯坦（Rubinstein）（1979）基于给定标的资产价格二元变动（价格上升或者下降）建立了一个期权定价模型。他们的二叉树期权定价模型是相关期权定价模型类中的第一个。该类模型在衍生品市场的发展中起到了重要的作用。这一个事实足以说明研究二项分布的原因，而且二项分布同样也可以应用于决策过程之中。

建立二项分布的基础是伯努利随机变量（Bernoulli random variable），它是用瑞士概率学家雅各布·伯努利（Jakob Bernoulli, 1654 - 1704）的名字来命名的。假设我们有一个每次产生两种结果之一的试验（一个重复发生的事件），该试验就被称为伯努利试验（Bernoulli trial）。如果我们当结果是成功时，令 Y 等于1，而当结果是失败时，令 $Y=0$ ，那么伯努利随机变量 Y 的概率函数将为

$$\begin{aligned} p(1) &= P(Y = 1) = p \\ p(0) &= P(Y = 0) = 1 - p \end{aligned}$$

式中， p 表示试验成功的概率。我们下面的例子是理解二叉树期权定价模型非常重要的第一步。

⊖ 参见 Hiller 和 Lieberman(2000)。由计算机最初产生的随机数一般为随机正整数，这些数被转化为近似的位于 0~1 之间的连续型均匀随机数。然后利用不同的方法将连续型均匀随机数产生其他分布（如正态分布）的随机观测值。我们将在蒙特卡罗模拟一节中进一步讨论随机观测值的产生方法。

例 5-2 将单期股票价格变动作为一个伯努利随机变量

假设我们以如下方式来描述股票价格的变动：今天股票价格为 S ，下一期的股票价格可能上升也可能下降。上升的概率为 p ，而下降的概率为 $1 - p$ 。因此，股票价格是一个具有成功（向上变动）概率等于 p 的伯努利随机变量。当股价上升时，最终价格变为 uS ，其中 u 等于 1 加上在股价上升情况下的收益率。例如，如果股票在上升的情况下获得 0.01 或者 1% 的收益，那么 $u = 1.01$ 。当股票下跌时，最终价格为 dS ，其中 d 等于 1 加上在股价下跌情况下的收益率。例如，如果股票在下跌情况下获得 -0.01 或者 -1% 的收益，那么 $d = 0.99$ 。图 5-1 展现了该股价动态变化模型的一个图表。

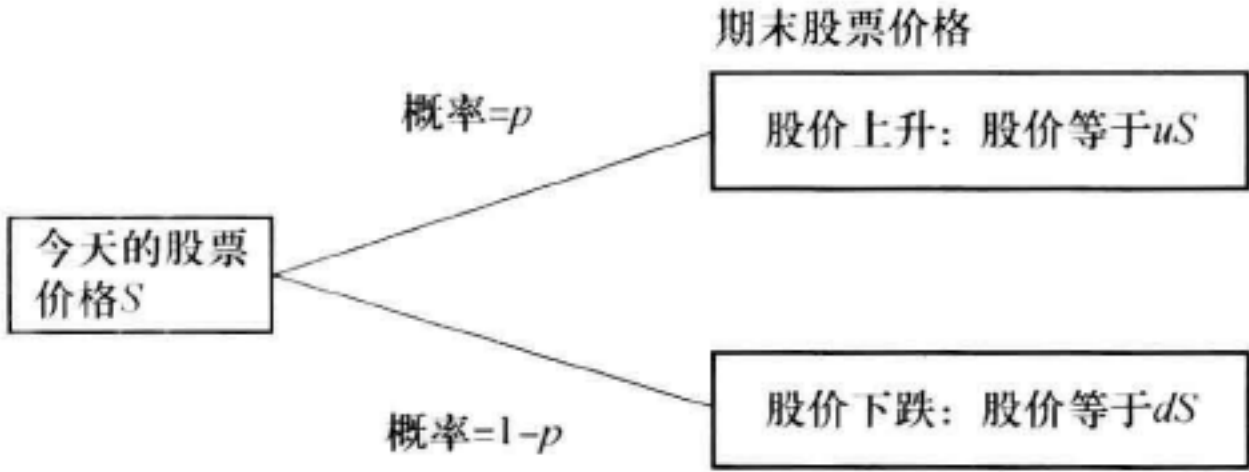


图 5-1 将单期股票价格作为单期伯努利随机变量

我们将在之后继续讨论上面的例子。在例 5-2 的股价变动模型中，在给定试验中成功和失败是与向上的变动和向下的变动分别相关的。在下面的例子中，成功指的是一个有利可图的交易，而失败指的是一个不能获利的交易。

例 5-3 一个评估大宗交易商的交易台（1）

作为一位从事股票交易的机构货币经理，你经常和一些大宗交易商进行交易。大宗交易（block）是指买卖量非常大以至于超出交易商网络或者股票交易市场能提供正常流动性的交易。已知你的公司对于某一类股票非常关注。当大宗交易商想要卖出他们认为你的公司可能感兴趣购买的大宗股票时，他们就打电话到你的交易台。你知道这些交易具有一定的风险。例如，如果交易商的客户（卖出股票的卖方）已经知道对于该股票不利的消息，或者是他通过所有渠道卖出的股票总数并没有真实地告知你，你可能面临立刻发生的损失。时不时地，你的公司会对与大宗交易商交易的业绩进行审计。你的公司会计算交易后的、经过市场风险调整的购买大宗交易商股票的美元收益率。在此基础上，你将每笔交易分为没有获利机会或者有利可图的两类。你已经在电子表格中汇总了与每位交易商交易的业绩表现，表 5-2 是摘自 2003 年 11 月的一部分信息。（交易商的姓名被编号为 BB001 和 BB002。）

表 5-2 大宗交易的获利和损失		
2003 年 11 月		
	获利的交易	损失的交易
BB001	3	9
BB002	5	3

将每笔交易看做一个伯努利试验。计算与两个大宗交易商在 2003 年 11 月交易的获利交易百分比。这些是对于每个交易商成功（获利）交易概率 p 的估计。

你公司记录了大宗交易商 BB001 的 $3 + 9 = 12$ 笔交易（行的总数）。因为 12 笔交易中的 3 笔是获利的，因此，获利交易的百分比为 $3/12 = 0.25$ 或者 25%。对于交易商 BB002，获利交易的百分比为 $5/8$ 或者 62.5%。一笔交易是一个伯努利试验，上面的计算提供了关于两个交易商获利交易（成功）概率的估计值。对于交易商 BB001，你的估计值为 $\hat{p} = 0.25$ ；对于交易商 BB002，你的估计值为 $\hat{p} = 0.625$ 。[⊖]

⊖ p 上的“帽子”表示这是 p 的一个估计量，是交易商获利交易的概率。

在 n 次伯努利试验中, 我们成功次数的可能值为 $0 \sim n$ 中的某个整数。如果一次试验的结果是随机的, 那么 n 次试验总成功的次数也是随机的。二项随机变量 (Binomial random variable) X 定义为 n 次伯努利试验中成功的次数。一个二项随机变量是伯努利随机变量 $Y_i, i=1, 2, \dots, n$ 之和:

$$X = Y_1 + Y_2 + \dots + Y_n$$

式中, Y_i 表示第 i 次试验的结果 (如果成功为 1, 失败则为 0)。我们知道一个伯努利随机变量是由参数 p 来定义的。试验次数 n 是一个二项随机变量的第 2 个参数。二项分布做出下面这些假设:

对于所有的试验, 成功概率 p 是常数。

试验间是相互独立的。

第 2 个假设很大程度地简化了问题。如果试验间是相关的, 那么计算 n 次试验中给定成功次数的概率将变得非常复杂。

在上面两个假设之下, 一个二项随机变量是完全由两个参数 n 和 p 来描述的。我们将其记为

$$X \sim B(n, p)$$

我们读做“ X 具有参数为 n 和 p 的二项分布。”你可以将一个伯努利随机变量看做 $n=1$ 时的一个二项随机变量: $Y \sim B(1, p)$ 。

现在我们可以来求解 n 次试验中成功 x 次的二项随机变量概率的一般表达式。如果每个区间段变得非常小, 我们可以将股票价格动态变化模型来刻画任何可能的股价运动。每个区间段是一个伯努利试验: 已知股价上升的概率为 p , 下降的概率为 $1-p$ 。成功是指股价向上变动的情况, x 是 n 个区间段 (n 次试验) 中股价向上移动或者成功的次数。在每个区间段股价变动间独立且 p 为常数的情况下, n 个区间段中股价向上变动的次数就是一个二项随机变量。我们现在就开始求解 $P(X=x)$ (二项随机变量的概率函数) 的表达式。

任何 n 个区间段的序列中正好有 x 个向上变动, 必然有 $n-x$ 个向下变动。我们有许多不同的方式来排列上升和下降的变动, 以产生一个上升总数为 x 的序列。在给定独立事件, 任何具有 x 个上升变动的序列一定以 $p^x(1-p)^{n-x}$ 的概率出现。现在我们需要用 x 个上升总数的不同序列的数目乘以这个概率。利用概率论中的一些概念一章中的基本计数结果, 我们有

$$\frac{n!}{(n-x)!x!}$$

种不同的在 n 次试验中出现 x 次上升变动 (或成功) 和 $n-x$ 次下降变动 (或失败) 的序列。回忆一下在概率论中的一些概念一章, n 的阶乘 ($n!$) 被定义为 $n(n-1)(n-2)\dots 1$ (并规定

$0!=1$)。例如 $5!=(5)(4)(3)(2)(1)=120$ 。组合公式 $\frac{n!}{(n-x)!x!}$ 可以记为

$$\binom{n}{x}$$

(读做“ n 中取 x 的组合”或者“ n 选 x ”)。例如, 在 3 期中, 有 3 个不同的具有两个向上变动的序列: UUD, UDU 和 DUU。我们可以证明:

$$\binom{3}{2} = \frac{3!}{(3-2)!2!} = \frac{(3)(2)(1)}{(1)(2)(1)} = 3$$

如果假设每个具有两个向上变动的序列具有 0.15 的概率, 那么在 3 期中具有两个向上变动的总概率为 $3 \times 0.15 = 0.45$ 。这个例子告诉我们, 对于服从 $B(n, p)$ 分布的 X , 在 n 次试验中产生 x 次成功的概率为:

$$p(x) = P(X = x) = \binom{n}{x} p^x (1 - p)^{n-x} = \frac{n!}{(n-x)!x!} p^x (1 - p)^{n-x} \tag{5-1}$$

一些分布是对称的，如正态分布，而另外一些分布是不对称或者偏向一侧的，例如对数正态分布。当试验成功的概率是 0.50 时，二项分布是对称的，但是在其他情况下，是不对称或者偏向一侧的。

通过一个对称的例子我们来对式 (5-1) (概率函数) 和 cdf 进行介绍。考虑一个服从 $B(n = 5, p = 0.50)$ 分布的随机变量。表 5-3 包含了对于该随机变量完整的描述。表 5-3 的第 4 列是第 2 列 (n 中取 x 的组合数) 乘以第 3 列中的 $p^x (1 - p)^{n-x}$ ；第 4 列给出了第 1 列每个向上变动数的发生概率。第 5 列对于第 4 列中的数进行了累积，其为累积分布函数。

表 5-3 二项概率分布， $p = 0.50, n = 5$

向上变动数 x (1)	达到向上变动数 x 的 可能情况数 (2)	每种情况的概率 (3)	x 的概率 $p(x)$ (4) = (2) $\times$ (3)	$F(x) = P(X \leq x)$ (5)
0	1	$0.50^0 (1 - 0.50)^5 = 0.03125$	0.03125	0.03125
1	5	$0.50^1 (1 - 0.50)^4 = 0.03125$	0.15625	0.18750
2	10	$0.50^2 (1 - 0.50)^3 = 0.03125$	0.31250	0.50000
3	10	$0.50^3 (1 - 0.50)^2 = 0.03125$	0.31250	0.81250
4	5	$0.50^4 (1 - 0.50)^1 = 0.03125$	0.15625	0.96875
5	1	$0.50^5 (1 - 0.50)^0 = 0.03125$	0.03125	1.00000

如果我们保持 $n = 5$ 不变，但是将试验成功发生的概率降低为 10%，那么会发生什么？对 $X = 0$ (没有向上变动) 的“每种情况的概率”将变为 59% 左右： $0.10^0 (1 - 0.10)^5 = 0.59049$ 。因为零次成功仍只有一种情况 (第 2 列)，故 $p(0) = 0.59049$ 。你可能想要检验一下，给定 $p = 0.10, P(X \leq 2) = 99.14\%$ ：两个或者更少变动的概率超过 99%。随机变量的概率一定大部分集中在 0、1 和 2 次向上变动的地方，出现更大向上变动结果的概率非常小。结果为 3 次或者更多向上变动的情况处于右侧的长尾部分，于是分布是右偏的。另一方面，如果我们设定 $p = 0.90$ ，我们将获得 $p = 0.10$ 分布的镜像分布。该分布是左偏的。

在理解了二项概率分布函数之后，我们可以继续之前大宗交易商的例子。

例 5-4 一个评估大宗交易商的交易平台 (2)

你现在想要评估例 5-3 中与大宗交易商交易的业绩表现。你从下面两个问题开始讨论：

- (1) 如果你支付给与你交易的交易商一个平均意义下的公平价格，那么该交易获利的概率将是多少？
 - (2) 这两个交易商是否达到了你在概率上的获利预期？
- 你也注意到交易商的业绩表现要在样本容量的基础上进行评估，而且你需要使用二项概率函数 (式 (5-1))。因此，你需要回答下列问题 (3) (参考例 5-3 中的数据)。
- (3) 在交易价格是公平的假设下，
 - (a) 计算与交易商 BB001 交易中有 3 个或者少于 3 个获利交易的概率。
 - (b) 计算与交易商 BB002 交易中有 5 个或者超过 5 个获利交易的概率。

解：

(1 和 2) 如果你交易的价格是公平的，那么 you 与交易商的交易中有 50% 是获利的。[⊖]而实

⊖ 当然，你需要对于交易后整体市场的走向进行调整 (牛市会对任何交易商的记录有利)，而且可能的话，还需要对于其他风险作出调整。这里假设我们已经完成了这些调整。

际与交易商 BB001 交易的获利交易率为 25%。因此, 交易商 BB001 没有达到你对于业绩表现的期望值。与交易商 BB002 的 62.5% 的交易是获利的, 因此超出了你的预期。

(3)

a. 对于交易商 BB001, 交易 (试验) 的次数为 $n=12$, 而其中的 3 次是获利的。你被要求计算 3 次或者小于 3 次的获利交易发生的概率, $F(3)=p(3)+p(2)+p(1)+p(0)$ 。

假设与 BB001 交易获利的概率为 $p=0.50$ 。那么在 $n=12$ 和 $p=0.50$ 情况下, 根据式 (5-1), 3 次获利交易发生的概率为

$$p(3) = \binom{n}{x} p^x (1-p)^{n-x} = \binom{12}{3} (0.50^3)(0.50^9) = \frac{12!}{(12-3)!3!} 0.50^{12} = 220(0.000\,244) = 0.053\,711$$

如果交易商 BB001 给予你公平价格的话, 12 次交易中正好有 3 次获利的概率为 5.4%。现在你需要计算其他的概率:

$$p(2) = \frac{12!}{(12-2)!2!} (0.50^2)(0.50^{10}) = 66(0.000\,244) = 0.016\,113$$

$$p(1) = \frac{12!}{(12-1)!1!} (0.50^1)(0.50^{11}) = 12(0.000\,244) = 0.002\,93$$

$$p(0) = \frac{12!}{(12-0)!0!} (0.50^0)(0.50^{12}) = 1(0.000\,244) = 0.000\,244$$

将所有的概率值相加, $F(3)=0.053\,711+0.016\,113+0.002\,93+0.000\,244=0.072\,998$ 或者 7.3%。如果交易商 BB001 给予你公平价格的话, 12 次交易中有 3 次或者少于 3 次获利的概率为 7.3%。

b. 对于交易商 BB002, 尽管有现在所显示的较好的交易结果, 你仍估计与该交易商交易获利的概率为 $p=0.50$ 。那么在 50% 的概率之下, 该问题为在 $n=8$ 次交易中有 5 次及 5 次以上获利的概率为: $1-F(4)=p(5)+p(6)+p(7)+p(8)$ 。你可以计算 $F(4)$, 并将其从 1 中减去, 或者你也可以直接计算 $p(5)+p(6)+p(7)+p(8)$ 。

如果交易商 BB002 给予你公平价格的话, 我们可以首先计算 8 次交易中正好有 5 次获利的概率:

$$p(5) = \binom{8}{5} (0.50^5)(0.50^3) = 56(0.003\,906) = 0.218\,75$$

该概率约为 21.9%。其他的概率为:

$$p(6) = 28(0.003\,906) = 0.109\,375$$

$$p(7) = 8(0.003\,906) = 0.031\,25$$

$$p(8) = 1(0.003\,906) = 0.003\,906$$

因此, $p(5)+p(6)+p(7)+p(8)=0.218\,75+0.109\,375+0.031\,25+0.003\,906=0.363\,281$ 或者 36.3%。[⊖] 一个 36.3% 的概率是非常大的; 尽管你在 2003 年 11 月与交易商 BB002 的交易是成功的, 但是与 BB002 以公平价格进行交易的概率仍应为 0.50。如果与 BB002 的一项交易从获利重新归为不获利的话, 则正好有一半的交易是获利的。因此, 你的交易台至少是以公平价格在与 BB002 交易; 在做出交易好于公平价格的结论之前, 你将可能需要获得更多的证据。

在这些交易中盈利和亏损的大小是另外一个值得考虑的重要因素。如果所有获利交易只获得小额的盈利, 而所有的不获利交易却损失惨重, 那么即使你大部分的交易是获利的, 你也将会在你的交易上输钱。

⊖ 在这个例子中, 所有进行的计算都是通过手工完成的, 但是二项概率和 cdf 函数同样也可以通过计算机电子表格程序来获得。

在下面的例子中，二项分布将帮助我们评估一个投资经理的业绩情况。

例 5-5 满足一个跟踪误差的要求

你现在正为一名养老金出资者工作。你已经指派了一名新的货币经理来管理 5 亿美元的跟踪 MSCI EAFE 指数（摩根士丹利欧澳亚远东指数，该指数是用来度量除美国和加拿大之外的发达市场股票表现的）的投资组合。在研究之后，你认为预计经理将跟踪误差保持在以季度度量的基准收益率上下 75 个基点（bps）的区间之内是合理的。[⊖]跟踪误差（tracking error）是投资组合的总收益率（总费用）减去基准指数（这里是 EAFE[⊗]）的总收益率。为了进一步量化期望值，如果总误差在 90% 的时间内处于上下 75 个基点的区间之内，我们的目标将得到满足。经理在 8 个季度中的 6 个满足该目标。当然，8 个季度中的 6 个是 75% 的成功率。但是经理的历史记录如何准确地与你 90% 的成功率和样本容量（8 个观测值）相关联？为了回答这个问题，你必须求解给定一个假定真实或者基准的成功率为 90%，业绩表现不比上述更好的概率。计算这个概率（用手算或者通过电子表格计算）。

具体地说，你要求解样本中的跟踪误差在 8 个季度中有 6 个或者少于 6 个落在上下 75 个基点的区间之中的概率。在 $n=8$ 和 $p=0.90$ 的条件下，这个概率为 $F(6)=p(6)+p(5)+p(4)+p(3)+p(2)+p(1)+p(0)$ 。首先有

$$p(6) = (8!/6!2!)(0.90^6)(0.10^2) = 28(0.005314) = 0.148803$$

接着求解其他概率：

$$p(5) = (8!/5!3!)(0.90^5)(0.10^3) = 56(0.00059) = 0.033067$$

$$p(4) = (8!/4!4!)(0.90^4)(0.10^4) = 70(0.000066) = 0.004593$$

$$p(3) = (8!/3!5!)(0.90^3)(0.10^5) = 56(0.000007) = 0.000408$$

$$p(2) = (8!/2!6!)(0.90^2)(0.10^6) = 28(0.000001) = 0.000023$$

$$p(1) = (8!/1!7!)(0.90^1)(0.10^7) = 8(0.00000009) = 0.00000072$$

$$p(0) = (8!/0!8!)(0.90^0)(0.10^8) = 1(0.00000001) = 0.00000001$$

将所有这些概率加总，你最终得到 $F(6) = 0.148803 + 0.033067 + 0.004593 + 0.000408 + 0.000023 + 0.00000072 + 0.00000001 = 0.186895$ 或者 18.7%。如果经理具有在 90% 的时间内满足你期望的能力，那么他将会以中等的 18.7% 概率在历史记录中表现出现在的记录（或者比这更差的记录）。

你可以利用其他评估概念，如跟踪风险，定义为跟踪误差的标准差，来评估经理的业绩表现。虽然上面的计算只是你得到关于经理业绩表现结论的一个方面。但是为了回答涉及成功率的问题，你需要具备利用二项分布进行计算的技巧。

经常在投资中用到两个描述变量为均值和方差（或者标准差，即方差的正的平方根）。[⊗]表 5-4 给出了二项随机变量的均值和方差的表达式。

⊖ 一个基点是 1% 的百分之一。

⊗ 一些从业者使用跟踪误差来描述我们之后所称做的跟踪风险，即投资组合和基准收益率之差的标准差。

⊙ 均值（或算术平均数）是一个分布或者数据集中所有值的和除以所有数据的总个数。方差是一个度量偏离均值偏差程度的指标。参见统计学概念和市场收益率一章中关于这些概念的详细论述。

表 5-4 二项随机变量的均值和方差

	均值	方差
伯努利分布 $B(1, p)$	p	$p(1 - p)$
二项分布 $B(n, p)$	np	$np(1 - p)$

因为单个伯努利随机变量 $Y \sim B(1, p)$ ，以概率 p 取 1，而以概率 $1 - p$ 取 0，其均值或者加权平均结果是 p ，其方差是 $p(1 - p)$ 。[⊖]一个普通的二项随机变量 $B(n, p)$ 是 n 个伯努利随机变量之和，因此，服从 $B(n, p)$ 分布的随机变量的均值为 np 。已知一个服从 $B(1, p)$ 分布的变量具有 $p(1 - p)$ 的方差，故一个服从 $B(n, p)$ 分布的随机变量的方差为其 n 倍，即 $np(1 - p)$ ，假定所有的试验间（伯努利随机变量间）是独立的。我们可以对于两个具有不同概率的二项随机变量计算如下：

随机变量	均值	方差
$B(n = 5, p = 0.50)$	$2.50 = 5(0.50)$	$1.25 = 5(0.50)(0.50)$
$B(n = 5, p = 0.10)$	$0.50 = 5(0.10)$	$0.45 = 5(0.10)(0.90)$

对于一个服从 $B(n = 5, p = 0.50)$ 分布的随机变量，期望的成功次数为 2.5，其标准差为 $1.118 = (1.25)^{1/2}$ ；对于一个服从 $B(n = 5, p = 0.10)$ 分布的随机变量，期望的成功次数为 0.50，其标准差为 $0.67 = (0.45)^{1/2}$ 。

例 5-6 一个债券投资组合中违约债券的期望数目

假设你是一位债券分析师，现在被要求估计一个被动的、由 25 个不同美国发行者发行的高收益债券投资组合中发行的债券在下一年违约的数目。该投资组合中债券的信用评级集中聚集在穆迪 B2 级/标准普尔 B 级周围，这意味着这些债券从它们支付利息和偿还本金的能力上来看，是处于投机级的。对于 B2/B 这个级别债券的年度估计违约率为 10.7%。

- (1) 假定用二项模型来对违约进行刻画，在下一年，该投资组合中违约债券的期望数目为多少？
- (2) 估计下一年违约债券数目的标准差。
- (3) 在这个背景下对二项概率模型的使用进行评论。

解：

- (1) 对于每个债券，我们可以定义一个伯努利随机变量，当债券在下一年中违约时，该随机变量取 1，不然则取 0。对于 25 只债券，下一年违约的期望值为 $np = 25(0.107) = 2.675$ 或者约为 3。
- (2) 方差为 $np(1 - p) = 25(0.107)(0.893) = 2.388\ 775$ 。标准差为 $(2.388\ 775)^{1/2} = 1.55$ 。因此一个关于违约期望数两倍标准差的置信区间将覆盖从 0 左右到 6 左右。
- (3) 二项模型的一个假设为试验间是独立的。在这个背景下，试验与每个债券发行者是否会在下一年违约相关。因为发行公司可能会面临同样的经济因素影响，试验间可能不是独立的。不过，为了对于违约数做出一个快速的估计，二项模型是恰当的。

⊖ 我们推导 $p(1 - p)$ 是伯努利随机变量的方差的过程如下：注意一个伯努利随机变量可以取两个可能值之一，1 或者 0： $\sigma^2(Y) = E[(Y - EY)^2] = E[(Y - p)^2] = (1 - p)^2p + (0 - p)^2(1 - p) = (1 - p)[(1 - p)p + p^2] = p(1 - p)$ 。

之前，我们考察了一个简单的单期股价变动的模型。现在我们将该描述股价变动的模型再扩展到连续三天。每天是一个独立的试验。股票以常数概率 p 向上变动（向上转变概率（up transition probability））；如果股价向上变动， u 等于 1 加上向上变动所带来的收益率。股票以常数概率 $1 - p$ 向下变动（向下转变概率（down transition probability））；如果股价向下变动， d 等于 1 加上向下变动所带来的收益率。我们在图 5-2 中画出股票变动的情况，其中我们将每个 $n = 3$ 的股票价格变动用表示时间的 t 来标注。图的形状告诉我们为什么它被称为二叉树（binomial tree）的原因。每个格中的值给出了连续的变动情况或者结果，树上的每个分支被称为一个结点（node）。例如在这个例子中，一个结点就是在一个具体时点股价可能的值。

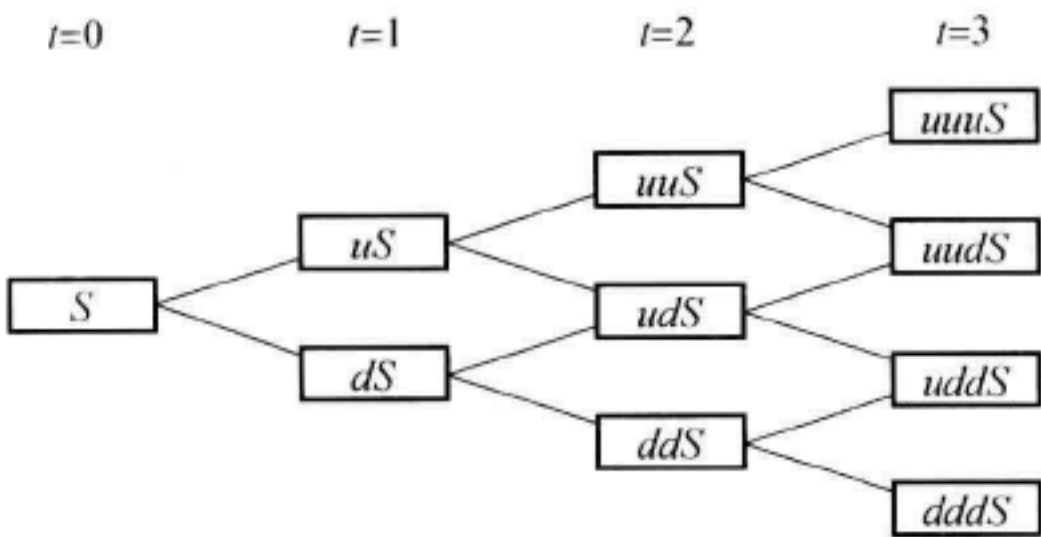


图 5-2 股价变动的二叉树模型

我们从树中看到在 $t = 3$ 上的股票价格有 4 个可能值： $uuuS$ ， $uudS$ ， $uddS$ 和 $dddS$ 。股票价格等于这 4 个值中的任一个概率是由二项分布所给出的。例如，有 3 个变动序列最终会导致了 $uudS$ 的结果：这些序列为 uud ， udu 和 duu 。这些序列在总的 3 次移动中有两次上升移动；组合公式也证明了在 3 个区间（试验）中有两次向上变动（成功）的种数为 $3!/(3-2)!2! = 3$ 。下面注意到这些序列中的每个 uud 、 udu 和 duu ，具有 $p^2(1 - p)$ 。所以 $P(S_3 = uudS) = 3p^2(1 - p)$ ，式中， S_3 表示 3 次移动后的股票价格。

在这个应用中的二项随机变量是向上变动的数目。最后股票价格的分布是关于初始股价、向上变动数目以及向上和向下变动总数目的一个函数。我们不能说股价本身是一个二项随机变量；但可以说它是二项随机变量的一个函数，同时也是 u 、 d 和初始价格的一个函数。这一补充实际上是解释为什么用这种方式对股价建模是有用的一个关键：它使我们选择这些参数来估计不同股票价格（利用多个时间段）。^①另一个可以被用来估计的分布是对数正态分布，这是一个我们将在后面讨论的重要的刻画股票价格的连续分布模型。它能够扩展得更加灵活。在上面显示的树中，转变概率在每个结点上是一致的：上升是 p ，下降是 $1 - p$ 。标准的公式描述了一个随着时间波动为常数的股价变化过程。然而，期权专家有时利用在不同结点上上升和下降概率不同的二叉树来对于随事件变动的波动率建模。

二叉树模型也提供了在任一结点检验一个条件或者选择权的可能性。这一选择的灵活性在诸如期权定价等投资应用中是非常有用的。考虑一个基于支付红利股票的美式期权。（回忆一下，一个美式看涨期权可以在到期日之前任何时候执行，在树上的任一结点执行。）在除息日之前，执行美式看涨期权，买入股票并收到红利是最优的。^②如果我们用二叉树对于股票价格建模，我们可以检验在每一结点上，是否执行期权是最优的。同样，如果我们知道看涨期权在 4 个 $t = 3$ 的终点上的价值，那么我们可以将它们的价值折现到前一期，即折现到 $t = 2$ 时刻，并计算出在 3 个结点处看涨期权的价值。继续逆向递归，我们可以计算出看涨期权在今天的价值。这种递归的处理方法很容易在计算机上进行编程。因此，二叉树模型可以对比美式看涨期权更加复杂的期权进行估值定价。^③

① 例如，我们可以将 20 天分成 100 个子区间，并小心地使用与之相匹配的 u 和 d 的值。
② 现金红利的发放意味着一个公司资产的减少。提前执行可能是最优的因为期权的执行价格一般并不会随现金红利发放的数额而减少，因此，现金红利对于美式看涨期权投资者所持有的头寸的影响是反向的。
③ 更多关于期权定价模型的信息参见 Chance（2003）。

5.3 连续型随机变量

在前一节中，我们对离散型随机变量（即随机变量的可能结果集会是可数的）进行了讨论。与之前的情况相反，连续型随机变量的可能结果是不可数的。如果 1.250 是一个连续型随机变量的可能值，那么我们不能说出下一个更大或者更小的可能值是多少。从技术上来说，一个连续型随机变量的可能结果的范围是一条实直线（所有在 $-\infty$ 与 $+\infty$ 之间的实数）或者是实直线的某个子集。

在这一节中，我们关注两个在投资工作中最重要的连续分布——正态分布和对数正态分布。与我们在离散分布中的处理一样，我们将通过均匀分布来引入该主题。

5.3.1 连续均匀分布

连续均匀分布是最简单的连续概率分布。均匀分布有两个主要用途。作为产生随机数技术的基础，均匀分布在蒙特卡罗模拟中起到重要的作用。而作为描述等可能结果的概率分布，均匀分布是一个恰当的用来表示某种不确定性的信念（即所有可能的结果是等可能的）的概率模型。

一个均匀随机变量的 pdf 是

$$f(x) = \begin{cases} \frac{1}{b-a} & \text{对于 } a < x < b \\ 0 & \text{其他} \end{cases}$$

例如，当 $a=0$ ， $b=8$ 时， $f(x) = 1/8$ 或者 0.125。我们将该密度画在图 5-3 上。该密度函数的图像是一条取值（高度）为 0.125 的水平直线。

一个边界为 $a=0$ 和 $b=8$ 的均匀随机变量小于或者等于 3 的概率是多少，或者 $F(3) = P(X \leq 3)$ ？当我们对于具有可能结果为 1, 2, ..., 8 的离散型均匀随机变量进行求解时，我们会将每个概率进行加总： $p(1) + p(2) + p(3) = 0.375$ 。而与此相反，对于一个连续型均匀随机变量或者任何连续型随机变量，取任一给定的常数的概率为 0。为了解释这一点，考虑一个 2.510 到 2.511 的小区间。因为这个区间

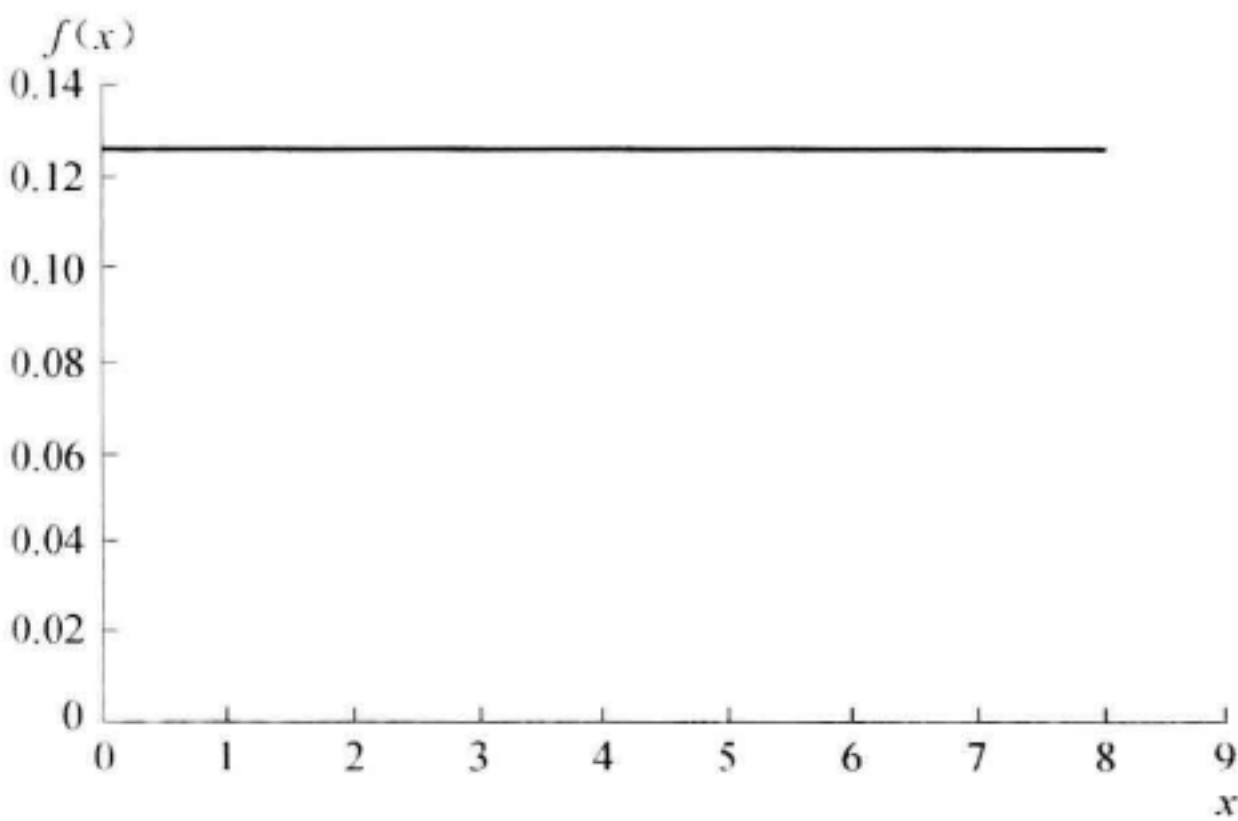


图 5-3 连续型均匀分布

内有无穷多个可能结果，如果每个可能值都拥有一个正的概率的话，那么仅在这个区间中所有值概率之和就将是无穷。所以为了计算 $F(3)$ ，我们要求解在 0~3 之间，描绘 pdf 的曲线图像下方的面积。在微积分中，这种运算称为对于概率密度函数 $f(x)$ 从 0 到 3 进行积分。在曲线下的面积为一个底为 $3-0=3$ 、高为 $1/8$ 的长方形。这个长方形的面积等于底乘以高： $3(1/8) = 3/8$ 或者 0.375。所以 $F(3) = 3/8$ 或者 0.375。

从 0 到 3 的这个区间是从边界 0 到 8 整个长度的 $3/8$ ，而 $F(3)$ 是总概率 1 的 $3/8$ 。cdf 中间的那行表达式刻画了这一关系。

$$F(x) = \begin{cases} 0 & \text{对于 } x \leq a \\ \frac{x-a}{b-a} & \text{对于 } a < x < b \\ 1 & \text{对于 } x \geq b \end{cases}$$

对于我们的问题来说，当 $x \leq 0$ 时，有 $F(x) = 0$ ，当 $0 < x < 8$ 时，有 $F(x) = x/8$ ，而当 $x \geq 8$ 时，有 $F(x) = 1$ 。我们在图 5-4 上画出了该 cdf。

与求解从 a 到 b 的一个 pdf, $f(x)$ 所表示的曲线下面积相对应的数学运算是从 a 到 b 对于函数 $f(x)$ 的积分：

$$P(a \leq x \leq b) = \int_a^b f(x) dx \quad (5-2)$$

式中， $\int dx$ 表示在小变化 dx 上的加总 $\int$ 符号，积分的上下界 (a 和 b) 可以为任意实数或者是 $-\infty$ 或者 $+\infty$ 。所有连续型随机变量可以用式 (5-2) 来进行计算。对于上面考虑的均匀分布的例子， $F(7)$ 在式 (5-2) 中的下限为 $a = 0$ ，而上限为

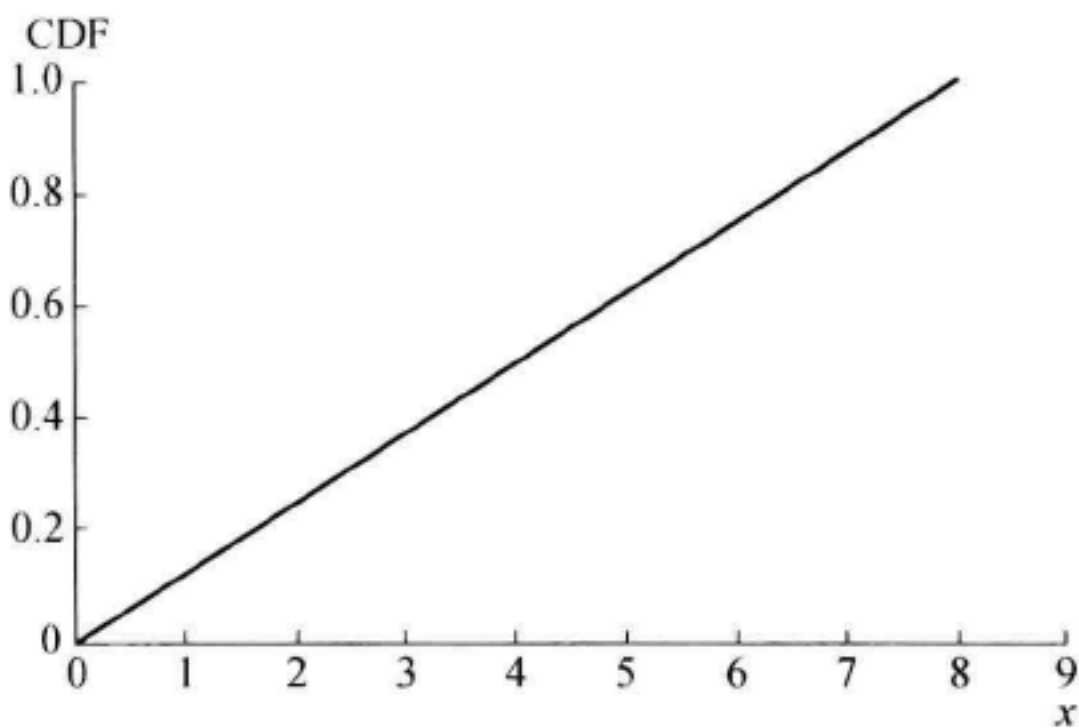


图 5-4 连续型均匀累积分布

$b = 7$ 。与均匀分布 cdf 相对应的积分简化为之前给出的 3 条直线。为利用式 (5-2) 来计算几乎所有的连续分布，包括正态分布和对数正态分布，我们要依靠电子表格中的函数、计算机程序或者是用来计算概率的数值表。这些工具使用数值方法来计算式 (5-2) 中的积分。

回忆一下，一个连续型随机变量等于任一固定值的概率为 0。这一事实对连续随机变量累积分布函数的计算产生这样一个重要的结果：对于任一连续型随机变量 X , $P(a \leq x \leq b) = P(a < x \leq b) = P(a \leq x < b) = P(a < x < b)$ ，因为在端点 a 和 b 的概率值为 0。对于离散型随机变量，这些等号并不一定是正确的，因为概率可能会累积在某点上。

例 5-7 一个信贷工具契约被违反的概率

你正在评估在商业周期低谷处一位低于投资级的借款者所发行的债券的信用情况。你有许多因素需要考虑，包括该公司银行借款工具的期限。签订一个银行借款工具的合同，如一个无担保的信贷，通常具有一些被称为贷款保证的契约条款。这些贷款保证的契约对于借款者能做什么有一定的限制。如果利息保障倍数（未计利息、税项、折旧及摊销前的利润与利息的比率（EBITDA/interest）（基于 4 个季度计算的 EBITDA）），低于 2.0 时，一家公司将会违反信贷工具中的贷款保证的契约。EBITDA 是未计利息、税项、折旧及摊销前的利润。[⊖] 是否符合贷款保证的契约将在当个季度末进行检查。如果贷款保证的契约被违反，银行可以要求立刻偿还该信贷工具出借的所有款项。这一行动将很可能造成该公司的流动性危机。你非常保守地预测所要求的利息为 25 000 000 美元。你估计 EBITDA 最低为 40 000 000 美元，最高为 60 000 000 美元。

回答下面两个问题（将所要求的利息支出处理为常数）：

⊖ 对于 EBITDA 使用方法（正确的使用与错误的使用）的深入讨论，参见 Moody's Investors Service Global Credit Research, *Putting EBITDA in Perspective* (June 2000)。

(1) 如果 EBITDA 的结果是等可能发生的, 那么 EBITDA/interest 低于 2.0, 即违反贷款保证的契约的概率是多少?

(2) 估计 EBITDA/interest 的均值和标准差。对于一个连续型均匀随机变量, 给出的均值为 $\mu = (a + b)/2$, 方差为 $\sigma^2 = (b - a)^2/12$ 。

解:

(1) EBITDA/interest 是一个连续型均匀随机变量, 因为所有结果是等可能的。该比率取值介于下限 $1.6 = (\$40\,000\,000)/(\$25\,000\,000)$ 到上限 $2.4 = (\$60\,000\,000)/(\$25\,000\,000)$ 之间。可能值的极差为 $2.4 - 1.6 = 0.8$ 。可能值低于 2.0 (引起违约的水平) 的比率是多少? 2.0 和 1.60 之间的距离为 0.40; 这个值 0.40 是总长度 0.80 的一半, 或者 $0.4/0.8 = 0.50$ 。因此, 违反贷款保证的契约的概率是 50%。

(2) 在解答 (1) 中, 我们求得 EBITDA/interest 的下限是 1.6, 记这个下限是 a 。而我们求得的上限是 2.4, 记这个上限是 b 。利用上面给出的公式,

$$\mu = (a + b)/2 = (1.6 + 2.4)/2 = 2.0$$

利率保障倍数的方差为:

$$\sigma^2 = (b - a)^2/12 = (2.4 - 1.6)^2/12 = 0.053\,333$$

标准差是方差正的平方根, $0.230\,940 = (0.053\,333)^{1/2}$ 。然而, 标准差作为均匀分布风险的度量并不是非常有用。位于偏离均值不同标准差区间的概率对于不同上、下限的设定是敏感的 (虽然切比雪夫不等式总是成立的)。^①这里, 均值 2.0 的 1 倍标准差区间是从 1.769 到 2.231, 涵盖了 $0.462/0.80 = 0.577\,5$ 或者 57.8% 的概率。两倍标准差区间从 1.538 到 2.462, 它超出了随机变量的上限和下限的范围。

5.3.2 正态分布

正态分布可能是在定量工作中使用最多的概率分布。它在现代投资组合理论和许多风险管理技术中发挥了重要的作用。因为它具有许多用途, 因此正态分布必须为投资从业者们所深入理解。

在统计推断和回归分析中正态分布的作用大部分是通过一个重要的结果 (中心极限定理) 而发挥出来的。中心极限定理认为大量的独立随机变量之和 (和均值) 近似服从正态分布。^②

法国数学家亚伯拉罕·棣莫弗 (Abraham de Moivre, 1667—1754) 在 1733 年建立中心极限定理的一个版本时引入了正态分布。如图 5-5 所示, 正态分布是对称且呈钟型的。

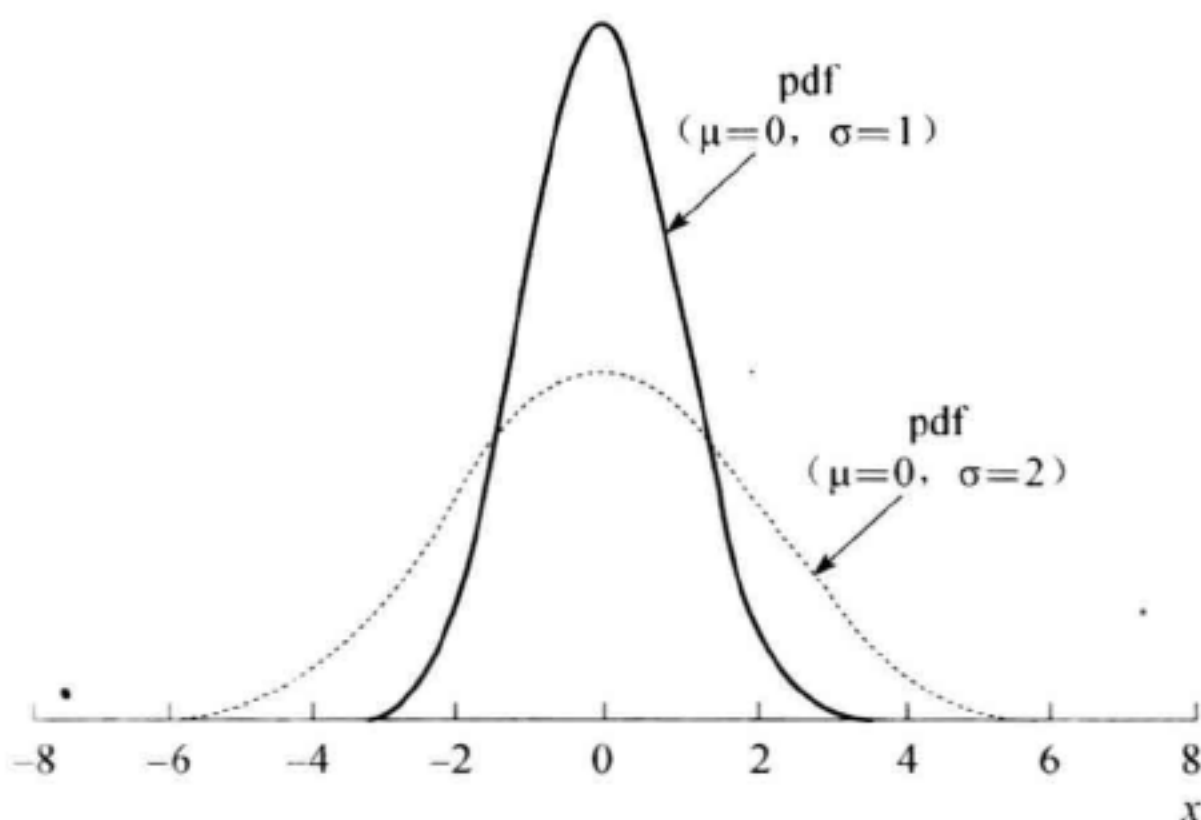


图 5-5 两个正态分布

① 切比雪夫不等式在统计学概念和市场收益率一章中进行了探讨。

② 中心极限定理将在抽样一章中深入讨论。

正态分布可能结果的范围是整条实直线：所有位于 $-\infty$ 到 $+\infty$ 之间的实数。钟形曲线的尾部向左和向右分别延伸到无穷。

正态分布的特征定义如下：

- 正态分布完全由两个参数——它的均值 μ 和方差 σ^2 所描述。我们将其记为 $X \sim N(\mu, \sigma^2)$ （读做“ X 服从一个均值为 μ ，方差 σ^2 的正态分布”）。我们也可以用均值和标准差 σ 来定义一个正态分布（这样做通常非常方便，因为 σ 度量的单位和 X 与 μ 是一样的）。结果，如果我们知道其均值与方差，就可以通过一个正态随机变量来回答任何概率的问题。
- 正态分布的偏度为 0（因为它是对称的）。正态分布的峰度（尖峰程度的度量）为 3；它的超额峰度（峰度 - 3.0）为 0。[Ⓐ] 作为对称的一个必然结果，均值、中位数和众数对于一个正态随机变量来说都是相等的。
- 两个或者多个正态随机变量的线性组合也是服从正态分布的。

这些要点是关于一个变量或单变量的正态分布：只有一个正态随机变量的分布。一个单变量分布（univariate distribution）只描述一个随机变量。而一个多变量（多元）分布（multivariate distribution）是对于一组相关随机变量的概率进行了描述。你在投资工作和阅读中将遇到多元正态分布（multivariate normal distribution），而你应该了解下面这些关于它的内容。

当我们有一组资产，我们可以对于每个资产的收益率分布分别建模，或者将所有资产看做一个组，对于该组的分布进行建模。“看做一个组”意味着我们会考虑所有收益率序列间统计上的相互关系。一个经常用来描述证券收益率的模型是多元正态分布。一个包含 n 种股票收益率的多元正态分布完全有下面 3 个参数所确定：

- 每个证券收益率的均值序列（一共 n 个均值）；
- 每个证券收益率的方差序列（一共 n 个方差）；
- 所有两两不同收益率间的相关系数序列：一共 $n(n-1)/2$ 个不同的相关系数。[Ⓑ]

需要对于相关系数进行定义是多元正态分布和单变量正态分布之间最主要的区别。

“假设收益率服从正态分布”的表述有时用来表示一个联合正态分布。例如对于一个由 30 只证券组成的投资组合，投资组合收益率是这 30 只证券收益率的加权平均值。一个加权平均值是一个线性组合。因此，如果单个证券收益率的收益率服从（联合）正态分布，那么投资组合的收益率也服从正态分布。回顾一下，为了定义投资组合收益率的正态分布，我们需要均值、方差和两两不同成分股票间的相关系数。

了解了这些概念之后，我们可以回到一个随机变量的正态分布。图 5-5 中所画的曲线是正态密度函数：

$$f(x) = \frac{1}{\sigma \sqrt{2\pi}} \exp\left(-\frac{(x - \mu)^2}{2\sigma^2}\right) \text{ 对于 } -\infty < x < +\infty \quad (5-3)$$

图 5-5 中画出的两个密度函数对应于均值 $\mu = 0$ 和标准差 $\sigma = 1$ 和 $\sigma = 2$ 。 $\mu = 0$ 和 $\sigma = 1$ 的正态

Ⓐ 如果我们一个大小为 n 的样本来自一个正态分布，我们可能想知道样本的偏度和峰度的可能变动。对于一个正态随机变量，样本偏度的标准差为 $6/n$ ，而样本峰度的标准差为 $24/n$ 。

Ⓑ 例如，一个包含两只股票的分布（一个二元正态分布）具有两个均值、两个方差和一个相关系数： $2(2-1)/2$ 。一个包含 30 只股票的分布具有 30 个均值、30 个方差和 435 个不同的相关系数： $30(30-1)/2$ 。美国陶氏化学与美国运通股票收益率间的相关系数和美国运通与美国陶氏化学股票收益率间的相关系数相同，因此，这两个相关系数只算为一个不同的相关系数。

密度函数被称为标准正态分布 (standard normal distribution)，或者单位正态分布 (unit normal distribution)。画出两个相同均值但不同标准差的正态分布有助于我们领会为什么标准差是正态分布离散程度的一个很好的度量： $\sigma=1$ 的正态分布的观测值比 $\sigma=2$ 的正态分布更集中于均值周围。

虽然表述并不是十分准确，但是正态分布可以被认为是刻画收益率的近似模型。几乎所有的正态随机变量的概率集中在以均值为中心的 3 倍标准差的区间内。对于现实中许多资产收益率均值和标准差，出现低于 -100% 结果的概率很小。在一个给定的应用中，该估计是否是一个实证的问题。例如，对于一个充分分散化的股票投资组合，季度和年度的持有期收益率比日和周的收益率更接近于正态分布。^①在大部分的股票收益率序列中，一个持续偏离正态分布的问题是其峰度比 3 要大，即厚尾问题。因此，当我们用正态分布估计股票收益率的分布时，我们应该注意到正态分布倾向于低估极端收益率的概率。^②期权收益率是偏向一侧的，因为正态分布是一个对称分布，我们在使用正态分布对于包含明显期权头寸的投资组合收益率进行建模时，应该尤其小心。

然而，作为收益率模型，正态分布不太适合作为资产价格模型。一个正态随机变量没有下界。这个特征在投资应用中有许多含义。一个资产价格只可能跌到 0，在这一点上资产没有任何价值。因此，从业者们一般不使用正态分布对资产价格进行建模。同样要注意：从任一水平的资产价格变动到 0 可以表述为 -100% 的收益率。因为正态分布低于 0 的部分延伸至负无穷，因此，它不能精确地用来为资产收益率建模。

在假定正态分布是一个利率变量恰当模型的前提下，我们可以利用下面的一些概率表述：

- 约 50% 的观测值落在 $\mu \pm (2/3) \sigma$ 的区间内。
- 约 68% 的观测值落在 $\mu \pm \sigma$ 的区间内。
- 约 95% 的观测值落在 $\mu \pm 2\sigma$ 的区间内。
- 约 99% 的观测值落在 $\mu \pm 3\sigma$ 的区间内。

一、二和三个标准差的区间在图 5-6 中给出了描述。该区间便于记忆，但是它们只是概率的一个近似表述。更精确的区间是 95% 的观测值落在 $\mu \pm 1.96\sigma$ 的区间内以及约 99% 的观测值落在 $\mu \pm 2.58\sigma$ 的区间内。

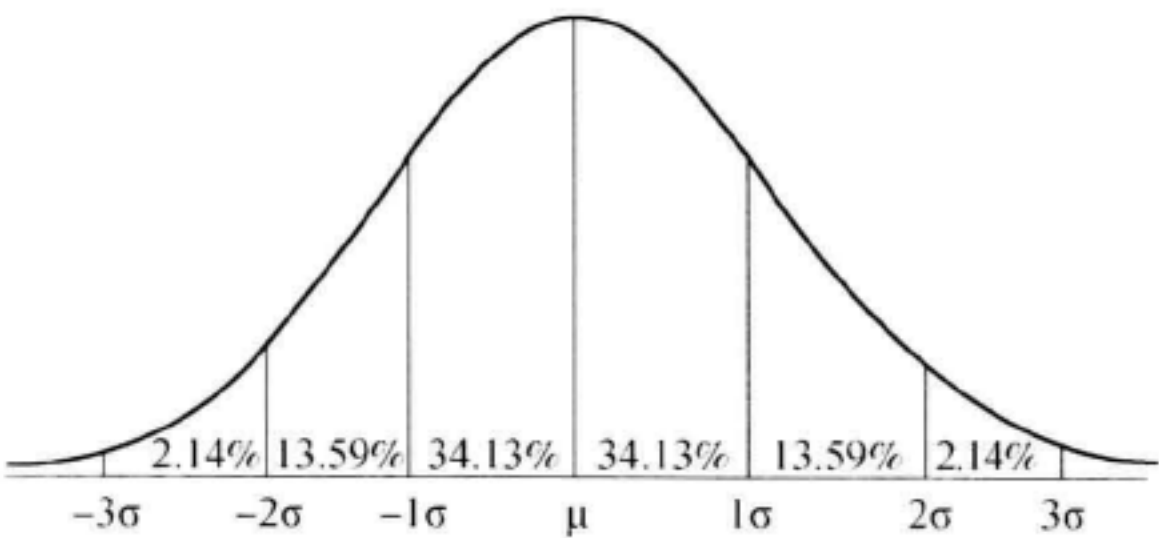


图 5-6 每单位的标准差

一般，我们不能观察到总体分布的均值或者标准差，所以我们需要估计它们。^③我们利用样本均值 $\bar{X}$ （有时记为 $\hat{\mu}$ ）来估计总体均值 μ ，用样本标准差 s （有时记为 $\hat{\sigma}$ ）来估计总体标准差 σ 。利用样本均值和标准差来分别估计总体均值和标准差，我们可以得到下面的关于服从正态分布随机变量的概率表述。（在这里，我们使用更精确标准差的数字来表述区间。）

正态随机变量 X 值的置信区间

- 我们预计 90% 的 X 的值落在区间 $\bar{X} - 1.65s$ 到 $\bar{X} + 1.65s$ 中。我们将这个区间称为 X 的

① 参见 Fama (1976) 和 Campbell、Lo 和 MacKinlay (1997)。

② 厚尾可以通过正态随机变量的混合或者是相对较小自由度的学生氏 t 分布来进行建模。参见 Kon (1984) 和 Campbell、Lo 和 MacKinlay (1997)。我们将在抽样和估计一章中讨论学生氏 t 分布。

③ 一个总体是一个特定组类中的所有元素，总体均值是由总体所计算出的算术平均值。一个样本是总体的一个子集，样本均值是由样本所计算出的算术平均值。更多关于这些概念的信息，参见统计学概念和市场收益率一章。

90%的置信区间。

- 我们预计95%的 X 的值落在区间 $\bar{X} - 1.96s$ 到 $\bar{X} + 1.96s$ 中。我们将这个区间称为 X 的95%的置信区间。
- 我们预计99%的 X 的值落在区间 $\bar{X} - 2.58s$ 到 $\bar{X} + 2.58s$ 中。我们将这个区间称为 X 的99%的置信区间。

例 5-8 一个由普通股组成的投资组合的概率 (1)

假如你管理着一个美国核心股票的投资组合。该投资组合对于标准普尔 500 指数是行业中性的 (它分配给各行业的权重近似与标准普尔 500 指数相同)。通过对于所持有股票均值收益率进行加权平均, 你预计该投资组合的收益率为 12%。你估计其年收益率的标准差为 22%, 接近于长期标准普尔 500 的收益率。对于该投资组合一年之后的收益率, 你要求回答下列问题:

- (1) 在收益率服从正态分布的假设下, 计算并解释该投资组合的一个标准差的置信区间。
- (2) 在收益率服从正态分布的假设下, 计算并解释该投资组合的 90% 的置信区间。
- (3) 在收益率服从正态分布的假设下, 计算并解释该投资组合的 95% 的置信区间。

解:

(1) 一个标准差的置信区间为 $\bar{X} \pm s$ 。已知 $\bar{X} = 12$ 以及 $s = 22\%$, 一个标准差区间的下界为 $-10\% = 12\% - 22\%$, 而上界为 $34\% = 12\% + 22\%$ 。因此, 这个区间是从 -10% 到 34% , 而你预计在正态假设下, 大约 68% 的投资组合收益率会落在该区间内。对于一个标准差置信区间的简单记号为 $[-10\%, 34\%]$ 。

(2) 在收益率服从正态分布的假设下, 一个 90% 的置信区间为从 $\bar{X} - 1.65s$ 到 $\bar{X} + 1.65s$ 。因此, 下界为 $-24.3\% = 12\% - 1.65(22\%)$, 而上界为 $48.3\% = 12\% + 1.65(22\%)$ 。简单地, 该区间为 $[-24.3\%, 48.3\%]$ 。

(3) 在收益率服从正态分布的假设下, 一个 95% 的置信区间为从 $\bar{X} - 1.96s$ 到 $\bar{X} + 1.96s$ 。因此, 下界为 $-31.12\% = 12\% - 1.96(22\%)$, 而上界为 $55.12\% = 12\% + 1.96(22\%)$ 。简单地来说, 该区间为 $[-31.12\%, 55.12\%]$ 。

95% 和 99% 的置信区间是两个在实践中最常用的置信区间。利用 2 而不是 1.96 作为乘数可以使更快地得到 95% 置信区间的估计, 因此这一方法也被经常应用。

解答 (3) 中下界 -31.12% 的计算说明了之前的一个结论: 对于许多实际的均值和标准差, 正态分布在左侧趋于 $-\infty$ 的事实可能并不值得批评。对于一个正态分布, 只有总概率的 2.5% 落在均值减去 1.96 倍标准差的左侧 (2.5% 落在均值加上 1.96 倍标准差的右侧)。图 5-7 通过展现总概率的 47.5% 位于均值加 (或减) 1.96 倍标准差之间来说明了这些概率。

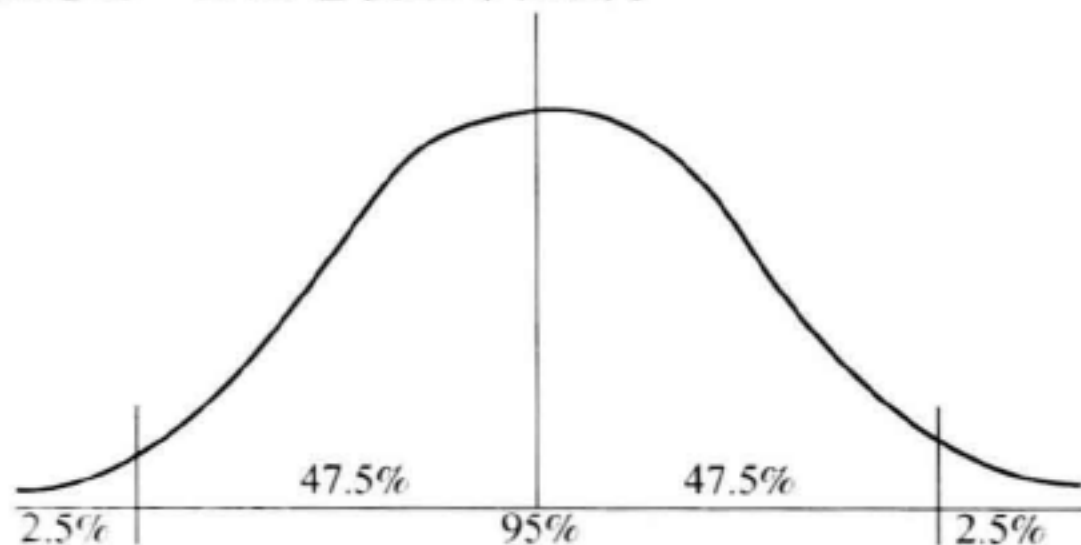


图 5-7 95% 置信区间的尾部概率

在对于置信区间的计算中, 我们设定了要求的置信水平, 并对于区间端点进行了计算。我们已经给出了重要的置信区间的计算公式, 但是也要对于其他的置信区间提出问题, 例如“包含投

投资组合 75% 收益率的置信区间有多宽?” 我们也会关注其他概率。例如, 我们可能会问, “股票指数的年收益率低于 1 年期国库券收益率的概率是多少?”

对于均值 (μ) 和方差 (σ^2) 的不同选择会产生许多不同的正态分布。我们可以利用任一正态分布回答上面所有的问题。例如, 电子表格具有任何均值和方差设定的正态 cdf 函数。然而, 为了效率的缘故, 我们将用一个一元正态分布来表示所有概率。标准正态分布 ($\mu = 0, \sigma^2 = 1$ 的正态分布) 可以充当该角色。

标准化 (standardizing) 一个随机变量 X 有两个步骤: 从 X 中将 X 的均值减去, 再除以 X 的标准差。如果我们有关于正态随机变量 X 的一系列观测值, 那么我们可以在列表中将每个观测值减去均值, 得到一系列偏离均值的偏差值, 再将每个偏差值除以标准差。所得到的结果是服从标准正态分布的随机变量 Z 。(Z 是传统用来表示标准正态分布的随机变量。) 如果我们有 $X \sim N(\mu, \sigma^2)$ (读作 “ X 服从参数为 μ 和 σ^2 的正态分布”), 我们利用下列公式将其标准化:

$$Z = (X - \mu) / \sigma \tag{5-4}$$

假设我们有 $\mu = 5, \sigma^2 = 1.5$ 的正态随机变量 X 。我们将用 $Z = (X - 5) / 1.5$ 对 X 标准化。例如, $X = 9.5$ 对应于标准化的值是 3, 计算为 $Z = (9.5 - 5) / 1.5 = 3$ 。我们将观测到小于等于 9.5 的 $X(X \sim N(5, 1.5^2))$ 的概率正好为小于等于 3 的 $Z(Z \sim N(0, 1))$ 的概率。我们能利用标准化的值和关于 Z 的概率数值表来回答所有关于 X 的概率问题。我们一般不知道总体的均值和标准差, 因此我们经常使用样本均值 $\bar{X}$ 来代替 μ , 用样本标准差 s 来代替 σ 。

标准正态分布也可以利用电子表格、统计计量软件和编程语言来进行计算。标准正态随机变量的累积分布表在本书的最后。表 5-5 列出了那张表的一部分。 $N(x)$ 是对于标准正态变量 cdf 的传统记号。[⊖]

表 5-5 对于 $x \geq 0, P(Z \leq x) = N(x)$ 或者对于 $z \geq 0, P(Z \leq z) = N(x)$

x 或 z	0	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.00	0.5000	0.5040	0.5080	0.5120	0.5160	0.5199	0.5239	0.5279	0.5319	0.5359
0.10	0.5398	0.5438	0.5478	0.5517	0.5557	0.5596	0.5636	0.5675	0.5714	0.5753
0.20	0.5793	0.5832	0.5871	0.5910	0.5948	0.5987	0.6026	0.6064	0.6103	0.6141
0.30	0.6179	0.6217	0.6255	0.6293	0.6331	0.6368	0.6406	0.6443	0.6480	0.6517
0.40	0.6554	0.6591	0.6628	0.6664	0.6700	0.6736	0.6772	0.6808	0.6844	0.6879
0.50	0.6915	0.6950	0.6985	0.7019	0.7054	0.7088	0.7123	0.7157	0.7190	0.7224

为了计算一个小于等于 0.24 的标准正态变量的概率, 我们可以寻找包含 0.20 的行, 然后找到 0.04 的列, 最后得到该栏中的数值 0.5948。因此, $P(Z \leq 0.24) = 0.5948$ 或者 59.48%

下面是在标准正态分布表中一些常被用来参考的数值。

- 第 90 个百分位点是 1.282: $P(Z \leq 1.282) = N(1.282) = 0.90$ 或者 90%, 而在右侧尾部还剩下 10% 的值。
- 第 95 个百分位点是 1.65: $P(Z \leq 1.65) = N(1.65) = 0.95$ 或者 95%, 而在右侧尾部还剩下 5% 的值。在计算百分位点时, 要注意区分单侧概率还是双侧概率。之前, 我们使用了对于 90% 置信区间的 1.65 倍标准差, 其中, 该置信区间的两侧分别各位 5%。这里我们

⊖ 另外一个经常看到的标准正态变量 cdf 的记号为 $\Phi(x)$ 。

使用 1.65, 因为我们在这里仅考虑位于一侧, 即右侧尾部的 5% 的值。

- 第 99 个百分位点是 2.327: $P(Z \leq 2.327) = N(2.327) = 0.99$ 或者 99%, 而在右侧尾部还剩下 1% 的值。

我们所给的正态 cdf 的表中包括 $x \leq 0$ 的概率。然而, 许多地方只给予 $x \geq 0$ 的数值表。如何使用该数值表来求解一个正态分布的概率? 因为正态分布的对称性, 我们可以利用标准正态分布变量 cdf 的数值表对于 $x \geq 0$, $P(Z \leq x) = N(x)$, 求得所有的概率值。下面的关系对于使用 $x \geq 0$ 的数值表和其他一些应用是有帮助的。

- 对于非负的数值 x , 使用表中的 $N(x)$ 。注意到对于 x 右侧的概率, 我们有 $P(Z \geq x) = 1.0 - N(x)$ 。
- 对于负的数值 $-x$, $N(-x) = 1.0 - N(x)$: 求解 $N(x)$ 并将其从 1 中减去。所有在正态曲线下 x 左侧的面积为 $N(x)$ 。余额 $1.0 - N(x)$ 是 x 右侧的概率和面积。利用均值周围的正态分布的对称性 x 右侧的概率和面积等于 $-x$ 左侧的概率和面积, $N(-x)$ 。
- 对于 $-x$ 右侧的概率, $P(Z \geq -x) = N(x)$ 。

例 5-9 一个由普通股组成的投资组合的概率 (2)

回忆一下, 在例子 5-8 中, 投资组合收益率的均值估计为 12%, 收益率的标准差估计为每年 22%。

利用这些估计值, 你想要计算下面这些概率, 假设用正态分布来描述收益率。(你可以使用正态分布概率表的一部分来回答这些问题。)

(1) 投资组合收益率超过 20% 的概率是多少?

(2) 投资组合收益率在 12% ~ 20% 之间的概率是多少? 换句话说, $P(12\% \leq \text{投资组合收益率} \leq 20\%) = ?$

(3) 你可以 5.5% 的收益率购买一只 1 年期的国库券。这一收益率是有效的 1 年期无风险收益率。你的投资组合收益率将小于等于无风险利率的概率是多少?

如果 X 是投资组合收益率, 标准化的投资组合收益率是 $Z = (X - \bar{X})/s = (X - 12\%)/22\%$ 。我们在该解答中使用这一表述。

解:

(1) 对于 $X = 20\%$, $Z = (20\% - 12\%)/22\% = 0.363636$ 。你想要求解 $P(Z > 0.363636)$ 。首先, 注意 $P(Z > x) = P(Z \geq x)$, 因为正态分布是一个连续型分布。回忆一下, $P(Z \geq x) = 1.0 - P(Z < x)$ 或者 $1 - N(x)$ 。将 0.363636 近似为 0.36, 根据数值表, 我们有 $N(0.36) = 0.6406$ 。因此, 如果你的正态假设是正确的话, $1 - 0.6406 = 0.3594$ 。投资组合收益率超过 20% 的概率约为 36%。

(2) $P(12\% \leq \text{投资组合收益率} \leq 20\%) = N(\text{对应于 } 20\% \text{ 的 } Z) - N(\text{对应于 } 12\% \text{ 的 } Z)$ 。对于第 1 项, $Z = (20\% - 12\%)/22\% = 0.36$ (大约), $N(0.36) = 0.6406$ (和解答 (1) 一样)。为了迅速获得第 2 项, 注意到均值是 12%, 对于正态分布 50% 的概率位于均值的任一边。因此, $N(\text{对应于 } 12\% \text{ 的 } Z)$ 必然等于 50%。因此 $P(12\% \leq \text{投资组合收益率} \leq 20\%) = 0.6406 - 0.50 = 0.1406$ 或者约为 14%。

(3) 如果 X 是投资组合收益率, 那么我们希望求解 $P(\text{投资组合收益率} \leq 5.5\%)$ 。这个问题比前两个问题更具挑战性, 但是当你已经学习了下面的解答之后, 你将有一个有用的计算其他差额概率的方法。

标准化投资组合收益率有 3 个步骤: 第 1 步, 将不等式每边中将投资组合的均值减去: $P(\text{投资组合收益率} - 12\% \leq 5.5\% - 12\%)$ 。第 2 步, 在不等式每边除以投资组合收益率的标准差: $P[(\text{投资组合收益率} - 12\%) / 22\% \leq (5.5\% - 12\%) / 22\%] = P(Z \leq -0.295455) = N(-0.295455)$ 。第 3 步, 等式左侧的部分是一个标准正态变量, 将其记为 Z 。正如我们上面所指出的, $N(-x) = 1 - N(x)$ 。将 -0.29545 近似为 -0.30 , 并把它用在摘要表中, 我们有 $N(-0.30) = 1 - N(0.30) = 1 - 0.6179 = 0.3821$, 近似为 38.21%。你的投资组合收益率低于 1 年期无风险收益率的概率约为 38%。

我们能通过从 5.5% 中减去投资收益率均值, 再除以投资组合收益率标准差, 并将结果 (-0.295455) 代入标准正态 cdf 来快速得到上面的答案。

5.3.3 正态分布的应用

现代投资组合理论 (modern portfolio theory, MPT) 经常使用的概念是投资机会的价值可以用收益率均值和方差来进行度量。在经济理论中, 当投资者是风险规避的时候, 或当他们选择投资组合以使得期望效用 (或满意程度) 最大化时, 或者当 (1) 收益率是服从正态分布的或 (2) 投资者是二次效用函数时,^① 均值 - 方差分析 (mean-variance analysis) 就是成立的。然而, 在假设 (1) 或者 (2) 其一违背时, 均值 - 方差分析仍可能是有用的——即其可能是近似成立的。因为从业者们更愿意处理诸如收益率之类的可观测变量, 而投资收益率至少近似服从正态分布的假设在大部分现代投资理论中发挥着重要的作用。

如果用标准差来对于高于和低于均值的波动率进行刻画, 均值 - 方差分析一般被认为是关于风险对称的。^② 另一种方法是只考虑下行风险。我们讨论这一种方法, 即安全第一法则, 因为它提供了一个正态分布理论在实际投资问题中应用的很好例子。安全第一法则 (safety-first rules) 关注于超亏风险 (shortfall risk), 该风险是指投资价值会低于某个时间区间内某个最低可接受水平的可能性。在固定养老金计划中的资产低于计划负债的风险就是一个超亏风险的例子。假设投资者将任何低于 R_L 水平的收益率视为不可接受的, 那么根据罗伊 (Roy) 的安全第一准则, 最优投资组合将最小化投资组合 R_p 低于临界值 R_L 的概率。^③ 用符号来表示, 投资者的目标是选择一个投资组合使得 $P(R_p < R_L)$ 最小化。当投资组合收益率服从正态分布时, 我们可以利用 R_L 位于期望投资组合收益率 $E(R_p)$ 下方的标准差的数值来计算 $P(R_p < R_L)$ 。 $E(R_p) - R_L$ 相对于标准差达到最大的投资组合能使 $P(R_p < R_L)$ 达到最小。因此, 如果收益率服从正态分布, 安全第一最优投资组合将最大化安全第一比率 (SFRatio):

$$\text{SFRatio} = [E(R_p) - R_L] / \sigma_p$$

$E(R_p) - R_L$ 这个量是平均收益率与损失水平之间的距离。将这个距离除以 σ_p 得到单位标准差的

① 效用函数是对于风险和收益偏好的数学表述。

② 我们将在投资组合概念一章中深入讨论均值 - 方差分析。

③ A. D. Roy (1952) 介绍了这个准则。

距离。运用罗伊准则选择投资组合有两个步骤（满足正态性假设）：[⊖]

- 1. 计算每个投资组合的 SFRatio。
- 2. 选择最高 SFRatio 的投资组合。

对于一个给定安全第一比率的投资组合，其收益率将低于 R_L 的概率为 $N(-\text{SFRatio})$ ，而且安全第一最优投资组合的该概率最低。例如，假设一个投资组合的临界收益率 R_L 为 2%。他被给予两个投资组合的选择。投资组合 1 具有 12% 的预期收益率和 15% 的标准差。投资组合 2 具有 14% 的预期收益率和 16% 的标准差。SFRatios 分别为投资组合 1 ($0.667 = (12 - 2)/15$)、投资组合 2 ($0.75 = (14 - 2)/16$)。对于较优的投资组合 2，假设投资组合收益率服从正态分布，投资组合收益率将低于 2% 的概率为 $N(-0.75) = 1 - N(0.75) = 1 - 0.7734 = 0.227$ 或者约为 23%。

你可能已经注意到 SFRatio 和夏普比率的相似性。如果我们将 SFRatio 中的临界水平 R_L 用无风险收益率 R_f 替代，那么 SFRatio 就变成了夏普比率。安全第一方法提供了一个看待夏普比率的新视角：当我们利用夏普比率对于投资组合进行评估时，具有最大夏普比率的投资组合就是使得投资组合收益率低于无风险利率概率最小的那个投资组合（给定正态假设的前提下）。

例 5-10 给客户提供的安全第一最优投资组合

你正在为一个客户提供的 800 000 美元进行资产配置。虽然她的投资目标是长期增长，但是在每年年末她可能要结算投资组合中的 30 000 美元以提供其所需的教育支出。如果这项需求出现，她可能会取出 30 000 美元而不影响最初的投资额 800 000 美元。表 5-6 列出了 3 种不同的分配方式。

表 5-6 三种分配方式的均值和标准差（以百分比为单位）

	A	B	C
收益率的均值	25	11	14
收益率的标准差	27	8	20

回答这些问题：（对于第 2 和第 3 题假设服从正态假定）

- (1) 给定客户不希望影响 800 000 美元本金的条件，损失水平 R_L 为多少？利用该损失水平回答第 2 个问题。
- (2) 根据安全第一准则，这 3 个分配方式哪个最优？
- (3) 安全第一最优投资组合的收益率小于损失水平的概率是多少？

解：

(1) 因为 $\$30\,000/\$800\,000$ 为 3.75%，对于任何小于 3.75% 的收益率，客户如果要取出 30 000 美元，都需要影响到本金。因此， $R_L = 3.75\%$ 。

(2) 为了决定 3 个分配方式中哪个是安全第一最优的投资组合，我们要选出最高的 $[E(R_p) - R_L]/\sigma_p$ 比率：

分配方式 A: $0.787\,037 = (25 - 3.75)/27$
分配方式 B: $0.906\,25 = (11 - 3.75)/8$

⊖ 如果存在一项资产提供所考察期限中的无风险收益率，而且如果 R_L 小于等于无风险收益率，那么完全投资于无风险资产就是最优的。在这个例子中持有无风险资产排除了没有达到临界值的可能性。

分配方式 C: $0.5125 = (14 - 3.75)/20$

分配方式 B 具有最大的比率 (0.906 25), 因此, 根据安全第一准则, 其是最佳的选择。

(3) 为了回答这个问题, 注意到 $P(R_B < 3.75) = N(-0.906\ 25)$ 。我们可以将 0.906 25 近似为 0.91, 从而利用标准正态分布的 cdf 数值表来获得答案。首先, 我们计算 $N(-0.91) = 1 - N(0.91) = 1 - 0.818\ 6 = 0.181\ 4$ 或者约 18.1%。如果不近似, 利用电子表格的函数, 也可以获得 -0.906 25 的标准正态 cdf。我们可以得到 18.24% 或者约为 18.2%。安全第一最优投资组合的收益率没有达到 3.75% 临界收益率水平的概率约为 18%。

这里有点问题值得我们注意。首先, 如果输入的数值有些许微小的变化, 我们可能会得到一个不同的排序。例如, 如果 B 的收益率均值是 10% 而不是 11%, 那么 A 就将优于 B。其次, 如果达到 3.75% 的临界收益率是必需的, 而不仅仅是一个希望, 那么 1 年后的 830 000 美元可以作为一项负债来建模。诸如现金流匹配的固定收益策略可以用来对冲或者免疫 830 000 美元的虚拟负债。

罗伊的安全第一规则是最早阐述损失风险的一个方法。标准均值方差投资组合选择过程也可以包含一个损失风险约束。^①

在除了罗伊安全第一准则之外的许多投资问题中, 我们也使用正态分布来估计概率。例如, 科尔布 (Kolb)、盖伊 (Gay) 和亨特 (Hunter) (1985) 推导出了一个基于标准正态分布的期货交易商会因为期货合约的损失而耗尽其现金流动性 (面临现金流动性问题) 的概率。标准正态分布发挥重要作用的另外一个领域是金融风险管理, 如投资银行、证券交易商和商业银行这些金融机构拥有正规的系统来度量和控制从交易头寸到企业总体风险等不同水平的风险。^② 金融风险管理中的两大支柱是在险价值 (Value at Risk, VAR) 和压力测试/情景分析。压力测试/情景分析 (stress testing/scenario analysis), 是 VAR 的一个补充, 它指的是一组用来估计在极端不利事件或者情景组合下损失情况的技术。在险价值是以货币度量的在一个特定区间 (例如, 一天、一个季度或者是一年) 中给定概率水平 (通常为 0.05 或者 0.01) 下的期望损失的最小值。假设我们设定为一天的时间区间和 0.05 的概率水平, 这可以被称为一个 95% 的一天的 VAR。^③ 如果这个 VAR 等于 5 000 000 欧元, 那么该投资组合有 0.05 的可能性在一天之内会损失 5 000 000 欧元或者更多 (假设我们的假定是正确的)。估计 VAR 的一个基本方法是方差 - 协方差或者分析方法, 它假设了收益率服从正态分布。更多关于 VAR 的内容, 可以参见强斯 (Chance) (2003)。

5.3.4 对数正态分布

对数正态分布与正态分布密切相关。该分布广泛地用于股票和其他资产价格概率分布的建模之中。例如, 对数正态分布就出现在布莱克 - 斯科尔斯 - 默顿期权定价模型中。布莱克 - 斯科尔斯 - 默顿期权定价模型假设期权的标的资产的价格服从对数正态分布。

一个随机变量 Y 服从对数正态分布, 那么其自然对数 $\ln Y$ 服从正态分布。反之也是正确的: 如果随机变量 Y 的自然对数 $\ln Y$ 服从正态分布, 那么 Y 服从对数正态分布。如果你把对数正态这

① 例如参见 Leibowitz 和 Henriksson (1989)。

② 金融风险 (Financial risk) 是指与资产价格和其他金融变量相关的风险。与其相对应的是与其他相关的非金融风险 (例如与经营和技术相关), 这种风险需要用其他不同的工具来进行管理。

③ 在 95% 的一天 VAR 中, 95% 是指对于 VAR 值的可信度, 其等于 $1 - 0.05$; 这是一个传统的表述 VAR 的方式。

个术语理解为“对数是正态的”，你将不会有任何困难来记住这个关系。

对数正态分布的两个最值得注意的结果是它以 0 为下界，而且偏向右侧（它具有右侧长尾）。这两个性质反映在图 5-8 的两条对数正态分布的概率密度函数（pdf）上。资产价格以 0 作为其下界。在实际应用中，对数正态分布被认为是一个对于许多金融资产价格分布有用的精确描述。而另一方面，正态分布经常作为收益率的一个很好的近似估计。由于这个原因，这两个分布对于金融从业者来说都非常重要。

与正态分布一样，对数正态分布完全由两个参数来刻画。而不同于其他我们所讨论过的分布，对数正态分布是由差分随机变量分布的

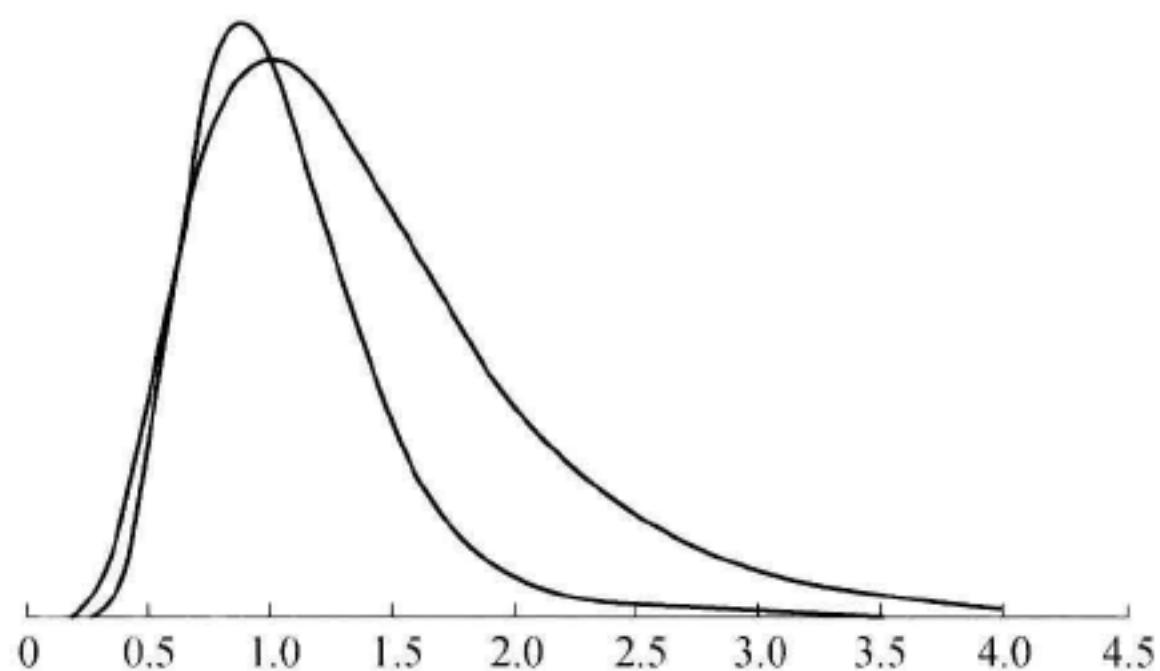


图 5-8 两个对数正态分布

参数来定义的。对数正态分布的两个参数分别是与其相联系的正态分布的均值和标准差（或方差）：给定 Y 服从对数正态分布，其参数为 $\ln Y$ 的均值和方差。记住，我们必须了解两组均值和标准差（或方差）：一组是相关的正态分布的均值和标准差（或方差）（这些是参数），另一组是对数正态变量本身的均值和标准差（或方差）。

对数正态变量自身均值和方差的表述是具有挑战性的。假设一个正态变量 X 具有期望值 μ 和方差 σ^2 。根据定义 $Y = \exp(X)$ 。我们知道 $\exp(X)$ 或者 e^X 的运算所隐含的是对数运算的逆运算。[Ⓐ]因为 $\ln Y = \ln[\exp(X)] = X$ 服从正态分布（我们假设 X 是正态变量），所以 Y 就是对数正态变量。那么 $Y = \exp(X)$ 的期望值是多少呢？一种猜测可能认为 Y 的期望值是 $\exp(\mu)$ 。而实际上，该期望值是 $\exp(\mu + 0.50\sigma^2)$ ，它要比 $\exp(\mu)$ 大一个因子 $\exp(0.50\sigma^2) > 1$ 。[Ⓑ]为了对这个概念有一些深入理解，我们考虑一下如果我们增加 σ^2 ，那么会发生什么。分布会伸展；该分布可以向上伸展，但是它不能向下伸展并超过 0。因此，该分布的中心被拉向右边，分布的均值增加了。[Ⓒ]

对数正态变量均值和方差的表述总结如下，其中 μ 和 σ^2 分别是对应正态分布的均值和方差（可以根据需要查看这些表述，而非记住它们）：

- 对数正态随机变量的均值 (μ_L) = $\exp(\mu + 0.50\sigma^2)$
- 对数正态随机变量的均值 (σ_L^2) = $\exp(2\mu + \sigma^2)[\exp(\sigma^2) - 1]$

我们现在来考察股票收益率和股票价格分布之间的关系。下面我们就来说明当一只股票的连续复利收益率服从正态分布时，那么未来的股票价格必然服从对数正态分布。[Ⓓ]并且我们会说明甚至当连续复利收益率不服从正态分布时，用对数正态分布也可以对股票价格很好地进行描述。这些结果提供了我们利用对数正态分布对于价格建模的理论基础。

为了概述下面的介绍，我们首先证明在未来 T 时刻的股票价格 S_T 等于当前股票价格 S_0 乘以 e

Ⓐ 量 $e \approx 2.7182818$ 。

Ⓑ 注意到 $\exp(0.50\sigma^2) > 1$ 因为 $\sigma^2 > 0$ 。

Ⓒ 该解释来源于 Luenberger (1998)。

Ⓓ 连续复利将时间处理为完全连续或者没有间断的。它与离散复利（将时间处理为向前的离散有限区间）是相对的。连续复利收益率是对于诸如布莱克-斯科尔斯-默顿期权定价模型所谓的连续时间金融模型中收益率的建模方式。对于复利的更多内容可以参见货币时间价值一章。

的 $r_{0,T}$ (0 到 T 时间区间的连续复利收益率) 次方; 这个关系可以表达为: $S_T = S_0 \exp(r_{0,T})$ 。于是, 我们可以证明我们可以把 $r_{0,T}$ 写成更短区间连续复利收益率的总和, 并且如果这些更短区间收益率服从正态分布的话, $r_{0,T}$ 也服从正态分布 (给定一定假设的条件下) 或者近似服从正态分布 (不作那些假设的条件下)。因为 S_T 与正态随机变量取对数后成比例, 故 S_T 服从对数正态分布。

为了提供一个讨论的框架, 假设我们有一组等间距的股价观测值: $S_0, S_1, S_2, \dots, S_T$ 。当前的股价为 S_0 , 它是一个已知的量, 并且是非随机的。而未来的价格 (如 S_1) 是随机变量。相对价格 (price relative) S_1/S_0 为期末价格 S_1 除以期初价格 S_0 ; 它等于 1 加上 $t=0$ 到 $t=1$ 区间的持有期股票收益率:

$$S_1/S_0 = 1 + R_{0,1}$$

例如, 如果 $S_0 = \$30$, 而 $S_1 = \$34.50$, 那么 $S_1/S_0 = \$34.50/\$30 = 1.15$ 。因此, $R_{0,1} = 0.15$ 或者 15%。一般地, 相对价格具有如下形式:

$$S_{t+1}/S_t = 1 + R_{t,t+1}$$

式中, $R_{t,t+1}$ 表示从 t 到 $t+1$ 的收益率。

一个重要的概念是与持有期收益率 (如 $R_{0,1}$) 相关的连续复利收益率。与持有期相关的连续复利收益率 (continuously compounded return) 为 1 加上持有期收益率的自然对数, 或者等价地, 等于期末价格除以期初价格 (相对价格) 的自然对数。^① 例如, 如果我们观察到一个一周持有期收益率为 0.04, 等价连续复利收益率 (称为一周连续复利收益率) 等于 $\ln(1.04) = 0.039221$, 1.00 欧元以 0.039221 的连续复利投资一周的收益为 1.04 欧元, 等价于 4% 的一周持有期收益率。从 t 至 $t+1$ 区间的连续复利收益率等于

$$r_{t,t+1} = \ln(S_{t+1}/S_t) = \ln(1 + R_{t,t+1}) \quad (5-5)$$

对于我们的例子, $r_{0,1} = \ln(S_1/S_0) = \ln(1 + R_{0,1}) = \ln(\$34.50/\$30) = \ln(1.15) = 0.139762$ 。因此, 13.98% 是区间 $t=0$ 到 $t=1$ 上的连续复利收益率。连续复利收益率小于相关的持有期收益率。如果我们的投资区间延长至 $t=0$ 到 $t=T$, 那么 T 时刻的连续复利收益率等于

$$r_{0,T} = \ln(S_T/S_0)$$

对于等式两边应用指数函数, 我们可以得到 $\exp(r_{0,T}) = \exp[\ln(S_T/S_0)] = S_T/S_0$, 所以

$$S_T = S_0 \exp(r_{0,T})$$

我们也可以将 S_T/S_0 表示成相对价格的乘积形式:

$$S_T/S_0 = (S_T/S_{T-1})(S_{T-1}/S_{T-2}) \cdots (S_1/S_0)$$

对这个等式两边取对数, 我们求得 T 时刻的连续复利收益率等于每期连续复利收益率之和:

$$r_{0,T} = r_{T-1,T} + r_{T-2,T-1} + \cdots + r_{0,1} \quad (5-6)$$

利用持有期收益率来求解 1 美元投资的最终价值涉及 (1 + 持有期收益率) 的连乘积, 而利用连续复利收益率涉及连加和。

在许多投资应用中的一个重要假设是收益率间是独立同分布的 (independently and identically distributed, IID)。独立性的假设来源于投资者不能利用历史收益率来预测未来收益率 (即弱式有效市场假设)。同分布的假设来源于平稳性的假设, 我们将在时间序列分析一章中介绍这个概念。^②

① 只在这章中, 我们使用小写的 r 来表示连续复利收益率。

② 平稳性表示收益率的均值和方差不会随着时间变化而改变。

假设一期连续复利收益率（如 $r_{0,1}$ ）是均值为 μ 方差为 σ^2 的 IID 随机变量（但是这里并没有假设其服从正态分布或者其他分布）。那么

$$E(r_{0,T}) = E(r_{T-1,T}) + E(r_{T-2,T-1}) + \cdots + E(r_{0,1}) = \mu T \tag{5-7}$$

（我们加总 μ 一共 T 次）以及

$$\sigma^2(r_{0,T}) = \sigma^2 T \tag{5-8}$$

（作为一个独立性假设的结果）。 T 个持有期连续复利收益率的方差是 T 乘以一期连续复利收益率；同样， $\sigma(r_{0,T}) = \sigma \sqrt{T}$ 。如果在式（5-6）右侧的一期连续复利收益率服从正态分布，那么 T 期持有期连续复利收益率 $r_{0,T}$ 也服从均值为 μT ，方差为 $\sigma^2 T$ 的正态分布。得到这样的关系是因为正态随机变量的线性组合仍是正态随机变量。但是即使一期连续复利收益率不服从正态分布，根据统计学中所知的中心极限定理的结果，它们的总和 $r_{0,T}$ 也是近似服从正态分布的。^① 现在将 $S_T = S_0 \exp(r_{0,T})$ 和 $Y = \exp(X)$ 相比较，其中 X 服从正态分布， Y 服从对数正态分布（如同我们上面所讨论的）。显然，我们可以将未来股票价格 S_T 以对数正态随机变量来进行建模，因为 $r_{0,T}$ 应该至少是近似正态的。这一对于服从正态分布收益率的假设是以对数正态分布作为股票价格和其他资产价格分布的模型理论为基础的。

正如前面所述，连续复利收益率在许多期权定价模型中发挥着重要的作用。波动率的估计是在应用诸如布莱克-斯科尔斯-默顿模型的期权定价模型时非常重要的。波动率（volatility）度量的是标的资产连续复利收益率的标准差。^② 在实际应用中，我们经常利用历史的连续复利日收益率序列来估计波动率。我们收集一组持有期日收益率，利用式（5-5）来将它们转化为连续复利日收益率。然后我们计算这组连续复利日收益率的标准差，并利用式（5-8）来将这个结果年化。^③（根据传统，波动率是以年化度量指标来表示的。）^④ 例 5-11 举例说明了如何估计米其林股票的波动率。

例 5-11 期权定价模型中所使用的波动率

假设你正在研究米其林公司（Euronext: MICP.PA）并对于其受到一系列国际事件影响的股票市场一周内的价格走势感兴趣。你决定使用波动率作为米其林股票在这一周内变动情况的度量指标。表 5-7 给出了这周的收盘价序列。

- 利用表 5-7 中的数据求解下列问题：
- （1）估计米其林股票的波动率（基于一年 250 天的年化波动率）。
 - （2）如果连续复利日收益率服从正态分布，判断米其林股价的概率分布。

表 5-7 米其林日收盘价

期	收盘价（欧元）
2003 年 3 月 31 日	25.20
2003 年 4 月 1 日	25.21
2003 年 4 月 2 日	25.52
2003 年 4 月 3 日	26.10
2003 年 4 月 4 日	26.14

资料来源：http://fr.finance.yahoo.com

① 在先前讨论正态分布时我们提到了中心极限定理。为了给出稍稍详细的表达，根据中心极限定理，无论这些随机变量服从什么分布，一组独立同分布有限方差的随机变量之和（均值也同样如此）将服从正态分布。我们将在抽样一章中对中心极限定理进行讨论。

② 波动率也被称为瞬时标准差，它被记为 σ 。标的资产，或者简单地称为标的，是指期权所对应的标的资产。关于这些概念更加详细的内容，可以参见 Chance（2003）。

③ 为了计算一组或者一个样本 n 个收益率的标准差，我们将每个收益率偏离平均收益率的平方偏差加总，然后再将总数除以 $n-1$ ，获得的结果就是样本方差，将样本方差求平方根得到样本标准差。要复习标准差的计算方法，可以参见统计学概念和市场收益率一章。

④ 年化经常是基于一年 250 天来进行计算的，这是对于市场开盘交易所估计的天数。250 天这个数字可能比 365 天更能获得一个好的波动率估计。因此，如果日波动率是 0.01，我们将会将波动率（基于年度基准）表述成 $0.01 \sqrt{250} = 0.1581$ 。

解:

(1) 首先, 利用式 (5-5) 来计算连续复利日收益率; 然后以常用的方式来求解它们的标准差。(在使用样本方差来获得样本标准差时, 利用样本容量减去 1 作为除数。)

$$\ln(25.21/25.20) = 0.000397, \ln(25.52/25.21) = 0.012222$$

$$\ln(26.10/25.52) = 0.022473, \ln(26.14/26.10) = 0.001531$$

$$\text{总和} = 0.036623, \text{均值} = 0.009156, \text{方差} = 0.000107$$

$$\text{标准差} = 0.010354$$

连续复利日收益率的标准差为 0.010354。式 (5-8) 显示 $\sigma^2(r_{0,T}) = \sigma^2 T$ 。在这个例子中, $\hat{\sigma}$ 是单期连续复利收益率的样本标准差。因此, $\hat{\sigma}$ 为 0.010354。我们想要将其年化, 因此时间长度 T 应该是 1 年。因为 $\hat{\sigma}$ 是以日为单位的, 我们将 T 设定为等于一年中交易天数 (250)。

我们求得那周米其林股票的年化波动率为 16.4%, 计算为 $0.010354 \sqrt{250} = 0.163711$ 。注意到样本均值 0.009156 是可能的连续复利单期或者日收益率均值 μ 的估计。样本均值可以利用式 (5-7) 转化为对于期望连续复利年收益率的估计: $\hat{\mu}T = 0.009156 \sqrt{250}$ (利用 250 以使得其与波动率的计算一致)。但是 4 个观测值对于预测期望收益率来说实在是太少了。日收益率的变动淹没了任何关于这么短序列的期望收益率的信息。

(2) 如果米其林股票连续复利日收益率服从正态分布, 那么米其林股价应该服从对数正态分布。

我们已经证明了给定一定假设条件下, 股价分布服从对数正态分布。那么, 如果 S_T 服从对数正态分布, 那么 S_T 的均值和方差是多少? 在这一节的前面部分, 我们给出了一个对数正态随机变量的均值和方差的重要表述。在这个重要的表述中, 我们所讨论的 $\hat{\mu}$ 和 $\hat{\sigma}^2$ 是与 S_T 的时间跨度相匹配的时间区间 T 的 (不是单期的) 连续复利收益率 (假设服从正态分布) 的均值和方差。^① 通过使用均值和方差 (或者标准差), 在这章的前面我们建立了一个包含正态分布随机变量特定比例观测值的一个区间。那些区间是关于均值对称的。那么我们能否对对数正态随机变量给出类似的对称区间表述呢? 不幸的是, 我们不能这么做。因为对数正态分布不是对称的, 这些区间比正态分布的情形要更加得复杂, 而我们也不在这里继续讨论这些专业的问题。^②

最后, 我们给出了不同时间区间的连续复利收益率的均值和方差之间的关系 (参见式 (5-7) 和式 (5-8)), 但是持有期收益率的均值和方差和连续复利收益率的均值和方差如何相联系? 例如作为分析师, 我们一般考虑的是持有期收益率而不是连续复利收益率, 我们可能在期权应用的时候希望将持有期收益率的均值和方差转化为连续复利收益率的均值和方差。为了推导这些转化公式 (以及从连续复利到持有期收益率的另一个方向的转化), 我们可以使用弗格森 (Ferguson, 1993) 中的表达式。

5.4 蒙特卡罗模拟

在了解了概率分布之后, 我们现在准备学习一个基于计算机的技术, 在这个技术中概率分布扮演着重要的角色。这个技术被称为蒙特卡罗模拟。金融中的蒙特卡罗模拟 (Monte Carlo simula-

① 例如, 均值的表达式为 $E(S_T) = S_0 \exp[E(r_{0,T}) + 0.5\sigma^2(r_{0,T})]$ 。

② 参见 Hull (2003) 对于对数正态置信区间的讨论。

tion) 涉及利用计算机来表现一个复杂的金融系统的运作过程。蒙特卡罗模拟的一个特征就是从特定的概率分布或者一组概率分布中产生大量的随机样本来反映系统中风险的影响。

蒙特卡罗模拟具有许多不同的应用。一个应用是在投资计划中。斯坦福大学的研究者山姆·萨维奇 (Sam Savage) 提供了下面这个形象的表述来反映这一应用: “在爬梯子前你会做的最后一件事是什么? 你会晃动它, 而这就是蒙特卡罗模拟。”^① 正如晃动梯子帮助我们判断爬上梯子的风险有多大, 蒙特卡罗模拟使我们在实际实施一个投资计划之前, 对于设定的投资方案进行试验。例如, 投资业绩表现可以通过与一个基准或者债务相比较而进行评估。养老金固定收益计划经常根据计划的负债来进行资产投资。养老金负债是一个复杂的随机过程。在蒙特卡罗资产-负债金融计划研究中, 养老金资产和负债的运行在给定资产如何进行投资、给定劳动力和其他变量的假设下, 在时间区间上被不断地进行模拟。在这里以及所有蒙特卡罗模拟中的一个重要设定为不同风险来源 (包括在这个例子中的利率和证券市场收益率) 的概率分布。对于计划融资状态下不同投资方案决策情况可以通过模拟时间序列来进行评估。试验可以对于不同的另一组假设重复进行。我们可以将下面的例 5-12 看做在这个标题下而得来的。在那个例子中, 市场收益率序列并不长得使研究者们回答股票市场择时的问题, 因此, 研究者们就模拟市场收益率来寻求他们问题的解答。

蒙特卡罗模拟也广泛应用于建立对于 VAR 的估计。在这个应用中, 我们模拟一定时间区间上投资组合的收益和损失的行为。在模拟中的重复试验 (每个试验涉及从一个概率分布中抽取一些随机观测值) 产生一个对于变化的投资组合值的频数分布。例如, 定义为最不利的 5% 的模拟变动的截断临界点是 95% VAR 的估计值。

作为一个非常重要的应用, 蒙特卡罗模拟是一个对于复杂证券估值的工具, 尤其是对于无法获得解析的定价公式表达式的欧式期权。^② 对于其他证券, 例如具有复杂的嵌入式期权的抵押贷款支持证券, 蒙特卡罗模拟也同样是一个重要的建模方法。

研究者们利用蒙特卡罗模拟来检验他们的模型和工具。一个特定的假设对于一个模型的表现到底有多重要? 因为我们在模拟过程中控制了假设条件, 因此我们可以通过蒙特卡罗模拟来运行模型以检验模型对于我们假设条件变化的敏感程度。

为了理解蒙特卡罗模拟技术, 让我们先将这个过程表示成一系列的步骤。^③ 为了介绍这些步骤, 我们举一个运用蒙特卡罗模拟来对于一种无法获得解析定价公式的期权 (基于一个股票的亚式看涨期权) 进行估值的例子来进行说明。一个亚式看涨期权 (Asian call option) 是指一个欧式期权, 其在到期日的价值等于到期日股票的价格和期权生存期内股票平均价格之差和 0 之间比较后的较大者。例如, 如果最终的股票价格为 34 美元, 而在期权生存期内的平均价格为 31 美元, 那么在到期日该期权的价值为 3 美元 ($\$34 - \$31 = \$3$ 和 $\$0$ 之间比较的较大者)。流程中的步骤 1~3 描述了对于模拟中要素的设定; 步骤 4~7 描述了模拟的实施。

1. 设定所关心的数值 (例如期权价值, 或者一个养老金计划的融资状态), 以一定的变量来表示。这个特定的变量或者一组变量可能是一个股票期权的股票价格, 养老金资产的市场价值, 或者其他与养老金计划中的养老金收益责任相关的变量。设定这些变量的初始值。

① 商业周刊 (Business Week), 2001 年 1 月 22 日。

② 一个欧式期权 (European-style option) 或者欧洲期权 (European option) 是指只能在到期日执行的期权。

③ 这些步骤可以看做为我们提供了一个蒙特卡罗模拟的总体流程, 而不是一个在许多不同应用中实施蒙特卡罗模拟的具体方案。

为了说明这些步骤，我们利用对于基于股票的亚式看涨期权的估值的例子来进行说明。我们用 C_{iT} 来表示期权在到期日 T 的价值。 C_{iT} 中的下标 i 说明这个 C_{iT} 的值是在第 i 次模拟试验（simulation trial）中得到的，每次模拟试验涉及抽取一组随机值（重复下面的步骤4）。

2. 确定时间的步长。将时间区间取为日历时间，然后将其分成一些子区间，比如说总共为 K 。日历时间除以子区间的数目 K 就是每个时间的增量 Δt 。
3. 确定驱动所设定变量风险因子的分布假设。例如，股票价格是亚式看涨期权的标的变量，因此我们需要一个模型来描述股票价格的变动过程。例如我们选择下面的模型来反映股价的变动，其中， Z_k 代表标准正态随机变量：

$$\Delta(\text{股票价格}) = (\mu \times \text{前一个时间的股票价格} \times \Delta t) + (\sigma \times \text{前一个时间的股票价格} \times Z_k)$$

以这种方式，我们来利用 Z_k 作为模拟中的风险因子。通过我们对于 μ 和 σ 的选取，我们控制了股票价格的分布。虽然这个例子只有一个风险因子，但是一个给定的模拟可能具有多个风险因子。

4. 利用一个计算机程序或者电子表格函数，对于每个风险因子抽取 K 个随机值。在我们的例子中，电子表格函数将会产生一个对于标准正态变量 Z_k 的 K 个抽取值： $Z_1, Z_2, \dots, Z_K$ 。
5. 利用步骤4中所产生的随机观测值计算所需变量的值。利用上面股票价格变动的模型，所得的结果是关于股价变化的 K 个观测值。所需的另外一个计算是将这些变动转化为 K 个股票价格（利用给定的初始股票价格）。此外，还有一个计算要来产生期权生存期内平均的股票价格（ K 个股票价格之和除以 K ）。
6. 计算所关心的数值。在我们的例子中，第一个计算是要得到亚式看涨期权在到期日的价值 C_{iT} 。第2个计算是要将这个最终值折现到当前以得到今天的看涨期权的价值 C_{i0} 。至此，我们完成了一次模拟试验。（ C_{i0} 的下标 i 表示第 i 次模拟试验，和之前 C_{iT} 中的一样。）在蒙特卡罗模拟中，一个运行表格是以在那个时点上模拟试验所得到的与所关心量（包括它们的均值和标准差）的分布相关的统计量来保存的。
7. 回到步骤4并重复进行直至达到所设定的实验次数 I 。最后，产生模拟所获得的统计量。对于我们例子中的重要数值是总的模拟试验次数中所获得的 C_{i0} 的均值。这个均值就是蒙特卡罗对于亚式看涨期权价值的一个估计值。

我们要设定多少次模拟试验？一般来说，我们需要用一个100的因子来增加试验次数以获得每个额外的准确数字。根据不同的问题，可能需要成千上万的试验来获得两个小数点的精确值（例如，为了获得期权的价值）。在给定当今计算机能力的条件下，执行许多模拟试验并不是一个问题。利用方差缩减方法可以降低所需的试验次数，而这个问题已经超出了本书所讨论的范围了。[⊖]

在我们例子中的步骤4，一个计算机函数产生一组基于一个标准正态随机变量的随机观测值。记得对于一个均匀分布，即所有可能产生的数值是等可能的。随机数生成器（random number generator）这个术语是指一个产生0到1之间均匀分布随机变量的算法。在计算机模拟的背景下，随机数（random number）这个术语是指从一个均匀分布中抽取一个观测值。[⊖]对于其他的分布，在

⊖ 对于这个问题以及其他的关于蒙特卡罗模拟的技术性问题，参见 Hillier 和 Lieberman（2000）。

⊖ 随机数生成器所产生的数值依赖于一个种子或者初始值。如果相同的种子被给予相同的生成器，它将会产生相同的数据序列。所有序列最终会相互重复。由于有这种可预测性，对于由随机数生成器所产生的数值的一个技术性正确名称应该称为伪随机数（Pseudo-random numbers）。伪随机数具有充分的随机性质以满足大部分实际应用的需要。

文中以“随机观测”这个术语来表达。

一个重要的事实是任何分布的随机观测值可以由0到1之间的均匀随机变量来产生。为了观察为什么是这样的，我们可以考虑产生随机观测值的逆变换方法。假设我们对于获得一个累积分布函数为 $F(x)$ 的随机变量 X 的随机观测值感兴趣。回忆一下， $F(x)$ 在 x 的值是0到1之间的某个数。假设一个随机变量的随机结果为3.21，而 $F(3.21) = 0.25$ 或者25%。定义一个 F 的反函数，称其为 F^{-1} ，它可以完成如下要求：将概率0.25代入 F^{-1} ，它将返回随机结果3.21。换句话说， $F^{-1}(0.25) = 3.21$ 。为了产生 X 的随机观测值，步骤为（1）利用随机数生成器产生一个0到1之间的均匀随机数 r 。（2）利用 $F^{-1}(r)$ 的值来得到一个 X 的随机观测值。随机观测值生成方式本身是一个研究的领域，我们这里所简要讨论的逆变换方法只是介绍了其中的一个观点。作为一个通才，你无需回答将随机数转化为随机观测值的具体技术性细节，但是你需要知道任何分布的随机观测值都可以通过一个均匀随机变量来产生。

在例5-12和5-13中，我们给出了一个应用蒙特卡罗模拟来回答投资实践中最为关心的问题（通过市场择时所获得的潜在收益）的例子。

例5-12 由市场择时所获得的潜在收益：蒙特卡罗模拟（1）

所有积极的投资者都想要实现优越的业绩表现。一个可能的优越业绩表现来源是投资者具有市场择时能力。一个投资者需要做出多么精确的预测才能成为一个具有获利能力的选择牛市和熊市时机的预测专家？在一个给定的精确水平之下，要获得多大的收益才能超过买入持有的策略？因为资产收益率具有变动性，我们需要大量的收益率数据来寻找统计上对于这些问题的可信的回答。因此，蔡（Chua）、伍德沃（Woodward）和土（To）（1987）选择蒙特卡罗模拟来回答由市场择时所获得的潜在收益的问题。他们是站在一个加拿大投资者的角度上进行研究的。

为了理解他们的研究，假设在年初，一个投资者预测该年将会出现一个牛市或者一个熊市。如果预测的是牛市，投资者将会把所有的钱投资于股票，并获得那一年的市场收益率。另一方面，如果预测的是熊市，投资者将会持有国库券，并获得国库券的收益率。基于上述事实，如果那一年的股票市场收益率 R_M 减去国库券收益率 R_F 为正，那么该市场就被归为牛市；反之，该市场就被归为熊市。一个市场择时者的投资结果可以与那些买入并持有的投资者进行比较。一位买入并持有的投资者每年将获得市场收益率。对于蔡等人的研究，一个令人感兴趣的量是从市场择时中获得的收益。他们定义这个量为市场择时者平均收益率减去买入并持有投资者的平均收益率。

为了模拟市场收益率，蔡等人产生了10 000个随机标准正态观测值 Z_t 。在研究的时点上，加拿大股票的历史年收益率均值为12.95，标准差为18.30。为了反映这些参数，模拟的市场收益率为 $R_{Mt} = 0.1830Z_t + 0.1295$ ， $t = 1, 2, \dots, 10\,000$ 。利用第2组10 000个随机标准正态观测值，加拿大国库券的历史收益率参数以及国库券和股票的历史相关系数，作者们产生了10 000个国库券的收益率。

投资者可能具有不同的预测牛市和熊市的能力。蔡等人将市场择时者们依照对于牛市和熊市的预测准确性进行分类。例如，牛市预测准确性为50%表示当择时者预测这年是牛市时，她

只在一半的时间中预测正确，这也意味着她没有预测能力。假设一个投资者具有 60% 的预测牛市的准确性和 80% 的预测熊市的准确性（一个 60~80 的择时者）。我们可以模拟该投资者是如何经营的。在产生第 1 个观测值 $R_{Mt} - R_{Ft}$ 之后，我们知道观测值是牛市还是熊市。如果观测值是牛市，那么 0.60（对于牛市预测的准确性）被与 0~1 之间的随机数相比。如果随机数小于 0.60（以 60% 的概率发生），那么市场择时者被认为具有准确预测牛市的能力，而她在第 1 个观测值时的收益率为市场收益率。如果随机数比 0.60 要大，那么市场择时者被认为预测错误而且预测了熊市；她在这个观测值时的收益率为无风险利率。以相同的方式，如果第 1 个观测值是熊市，该择时者具有 80% 的机会基于抽取的随机数预测正确熊市的出现。在任何一种情况下，她的收益率将与市场收益率相比较以记录其相对于买入并持有策略的收益情况。这个过程是一次模拟。模拟所获得择时者的平均收益率为在所有模拟试验中择时者所获得的平均收益率。

为了增进我们对于该过程的理解，考虑一个虚拟的 4 次对于 60-80 择时者（她不断重复地具有 60% 预测牛市的精确性和 80% 的预测熊市的精确性）的蒙特卡罗模拟。表 5-8 给出了模拟所获得的数据。让我们看一下试验 1 和试验 2。在试验 1 中，第 1 个抽取的随机数导致市场收益率为 0.121。因为市场收益率 0.121 超过国库券收益率 0.050，因此我们有一个牛市。我们产生一个随机数 0.531，我们用它与择时者牛市预测准确性 0.60 进行比较。因为 0.531 小于 0.60，所以择时者被认为做出了一个准确的牛市预测，并且她已经投资于股票。因此，该择时者在这次试验中获得了 0.121 的收益率。在第 2 次试验中，我们观测到另外一个牛市，但是因为随机数 0.725 大于 0.60，所以择时者被认为预测错误，她预测出现熊市了；因此，择时者只获得了国库券收益率 0.081，而不是更高的股票市场收益率。

表 5-8 虚拟的对于一个 60-80 市场择时者的模拟

试验	在抽取 Z_t 和国库券收益率之后			模拟结果		
	R_{Mt}	R_{Ft}	牛市还是熊市?	X 的值	择时者的预测是否正确?	择时者所获得的收益率
1	0.121	0.050	牛市	0.531	正确	0.121
2	0.092	0.081	牛市	0.725	错误	0.081
3	-0.020	0.034	熊市	0.786	正确	0.034
4	0.052	0.055	A	0.901	B	C
						$\bar{R} = D$

注： $\bar{R}$ 是择时者在 4 次模拟试验中所获得的平均收益率。

利用表 5-8 的数据，确定 A、B、C 和 D 的值。

解：A 的值是熊市，因为试验 4 中的股票市场收益率小于国库券收益率。B 的值是错误。因为我们观测到一个熊市，我们将随机数 0.901 和择时者预测熊市的准确性 0.80 相比较。因为 0.901 大于 0.8，所以择时者被认为预测错误。C 的值是股票市场的收益率 0.052，因为择时者预测错误，并且投资于股票市场获得了 0.052 的收益率而不是更高的国库券收益率 0.055。D 的值为 $\bar{R} = (0.121 + 0.081 + 0.034 + 0.052) / 4 = 0.288 / 4 = 0.072$ 。注意我们可以计算除均值以外的其他统计量，如择时者在这 4 次模拟试验中所获得收益率的标准差。

例 5-13 由市场择时所获得的潜在收益：蒙特卡罗模拟（2）

在讨论了蔡等人研究的计划和举例介绍了一个虚拟的 4 次试验蒙特卡罗模拟的方法之后，我们要总结一下这项研究的结论。

在例 5-12 中的虚拟模拟具有 4 次试验，这个数目太少了以致于无法达到统计上的精确结论。蔡（Chua）等人的模拟包含了 10 000 次试验。蔡等人设定牛市和熊市的预测能力水平为 50%，60%，70%，80%，90% 和 100%。表 5-9 给出了在没有交易成本情况下（他们也检验了交易成本存在的情况）的一个他们模拟结果的很小部分的摘要。从行上来看，具有 60% 牛市预测精确性和 80% 熊市预测精确性的择时者具有从市场择时中获得的每年平均 -1.12% 的年收益率。平均来看，买入并持有的投资者比具有这种能力的投资者多获取了 1.12 个百分点的收益。然而，在模拟试验之间的收益情况具有很大的变动性：收益的标准差为 14.77%，同样在许多试验中（但是不是在平均意义下），收益为正。第 3 行（获利/损失）是在股票和国库券中转换而获利的次数和损失次数之间的比率。这个比率对于 60-80 择时者来说是一个较有利的 1.207 0。（然而，当考虑交易成本时，更少的转换是获利的：对于 60-80 择时者的获利/损失比率为 0.583 2。）

表 5-9 由市场择时所获得的收益（不考虑交易成本）

牛市预测		熊市预测的准确性 (%)						
的准确性 (%)		50	60	70	80	90	100	
60	均值 (%)		-2.50	-1.99	-1.57	-1.12	-0.68	-0.22
	标准差 (%)		13.65	14.11	14.45	14.77	15.08	15.42
	获利/损失		0.741 8	0.906 2	1.050 3	1.207 0	1.349 6	1.498 6

资料来源：蔡，伍德沃和土（1987），表 II（摘要）

作者得到的结论是在牛市中不投资于市场的成本是巨大的。因为一个买入并持有的投资者从不错过任何一个牛市，因此，她 100% 能预测牛市的出现（以 0% 预测熊市作为代价）。给定他们的定义和假设，作者还认为成功的市场择时需要至少 80% 的对于牛市和熊市的准确性。市场择时仍是一个人们感兴趣的研究领域，其中还存在着许多其他的观点。然而，这个例子告诉了我们，蒙特卡罗模拟是如何被用来回答这些重要的投资问题的。

分析师要选择蒙特卡罗模拟中所使用的概率分布。相反，历史模拟（historical simulation）是通过从一个收益率（或者其他所关注的变量）的历史记录中抽样而对于一个过程进行的模拟。关于历史模拟的概念（也被称为后向模拟（back simulation））是指历史记录提供大部分直接的关于分布的证据（即把过去应用于未来）。例如，回到上面蒙特卡罗模拟概述中步骤 2，并假设时间的增量是一天。另外还假设我们的模拟基于最近 5 年的股票日收益率的记录。在一种类型的历史模拟中，我们从记录中随机抽取 K 个收益率来产生一个模拟试验。然后我们将获得的观测值放入样本中，在下一次的试验中，我们再做一次有放回的随机抽样。这个模拟的结果直接反映了数据的频数。这种方法的缺点在于任何不在该选择时间区间中所存在的风险将不会反映到模拟之中。与蒙特卡罗模拟相比，历史模拟并不适用于假设分析。但是，历史模拟也是一种已经建立的可供选择的模拟方法。

蒙特卡罗模拟是对于分析方法的一种补充。它仅仅提供了统计上的估计，而不是精确的结果。分析方法，如果可以获得的话，能使我们了解到更加深刻的因果关系。例如，布莱克-斯科尔斯-默顿期权定价模型对于欧式看涨期权的估值就是一种分析方法，它是以公式来表达的。在对于该类期权定价时，分析方法是比蒙特卡罗模拟更加有效的一种方法。作为一个分析表达式，布莱克-斯科尔斯-默顿模型帮助分析师快速地度量出期权价值对于当前股票价值和其他决定期权价格变量变动的敏感性。相反，蒙特卡罗模拟并不直接提供这种精确的结果。然而，只有一些类型的期权可以用分析表达式来定价。随着金融产品创新的不断向前推进，蒙特卡罗模拟的应用领域将会不断地扩大。

抽样和估计

6.1 引言

每天，我们都会观察到全球各股票市场指数的最高价、最低价和收盘价。诸如标准普尔 500 指数和日经 - 道琼斯平均工业指数这些指数代表的是市场所有股票的一个样本。虽然标准普尔 500 和日经指数并不能代表美国或者日本市场中所有股票的总体，但是我们常将它们视为反映整个总体表现的有效指标。作为分析师，我们习惯于利用这些样本信息来评估全球不同市场的表现情况。然而，我们通过样本信息所计算出的任何统计量，只是相对应总体参数的一个估计。所以，一个样本是总体的一个子集，我们要通过对该子集的研究来推断出总体自身的结论。

本章将探讨如何抽样以及如何利用样本信息来估计总体参数的问题。在下面一节中，我们将先讨论抽样（sampling）——获得一个样本的过程。在投资活动中，我们经常会利用均值作为随机变量（如收益率和每股盈利）集中趋势的度量指标。即使当随机变量的概率分布未知时，我们也能利用中心极限定理对于总体均值的概率情况进行表述。在第 3 节中，我们将讨论和介绍这个重要结果。在这些讨论之后，我们将转向参数估计。参数估计所寻求的是对于问题“参数值是多少？”的准确回答。

中心极限定理和参数估计是本章所介绍方法中的两个核心内容。在投资活动中，我们会将这些方法和其他统计技术应用到金融数据之上；我们经常为了确定在投资中哪些因素起作用、哪些因素不起作用，而对于这些统计结果进行解释。我们在本章的最后，将讨论如何对金融数据所得到的统计结果进行解释以及在这个解释过程中可能遇到的一些陷阱。

6.2 抽样

在这节中，我们会介绍通过样本（总体的部分元素）获得关于一个总体（一个特定组别的所有元素）信息的不同方法。我们试图获得的一个总体的信息经常是关于一个参数（parameter）的值，即一个由总体中数据计算得来的或者用于描述总体的数值。当我们利用一个样本来估计一个

参数时，我们会利用样本统计量（简称为统计量）。一个统计量（statistic）是一个由样本数据计算得来的或者用来描述样本的数值。

我们会出于下面两个原因之一而进行抽样。在有些情况中，我们可能无法检验总体中的每个元素。而在另一些情况中，检验总体的每个元素将是不经济的。因此，节省时间和金钱是造成分析师使用抽样方法来回答关于总体相关问题的两个主要原因。在这一节中，我们将讨论随机抽样的两种方法：简单随机抽样和分层随机抽样，然后我们会定义和解释分析师所使用的两种类型的数据：横截面数据和时间序列数据。

6.2.1 简单随机抽样

假设一个电信设备行业的分析师想要了解大部分用户未来一年将在设备上平均花费多少。一种方法是去调查电信设备用户的总体并询问他们具体的购买计划。用统计的术语来说，用户计划支出总体的特征通常可以表示为诸如均值和方差的描述性度量指标。然而，对于所有公司进行调查将花费非常多的时间和金钱。

另外一种方法是分析师可以收集公司的代表性样本，并调查它们在未来的电信设备支出。在这个例子中，分析师将计算样本均值支出 $\bar{X}$ ，这是一个统计量。这种方法比调查整个总体更具有优势，因为它可以以最小的成本迅速地完成任务。

然而抽样会引入误差。误差的产生是因为我们并没有调查所有总体中的公司。决定抽样的分析师会以产生抽样误差的代价换取时间和金钱的节省。

当一位分析师选择采用抽样方法时，他必须首先构建一个抽样计划。一个抽样计划（sampling plan）是一套用来选择样本的规则。我们可以从中得到统计上对于总体坚实结论的基本样本类型是简单随机样本（simple random sample，简称为随机样本）。

- **简单随机样本的定义。**一个简单随机样本是指一个较大总体的一个子集，它是通过等概率地选取总体中的每个元素的方式得到的。

以满足简单随机样本定义的抽取样本的过程被称为简单随机抽样（simple random sampling）。那么简单随机抽样是如何进行的呢？我们需要一个在选择样本过程中保证随机性的方法（不需要任何的模式）。对于一个有限（限定的）总体，最常用的得到一个随机样本的方法是利用随机数（保证具有随机性的一些数字）。首先，我们将总体中的元素按先后顺序进行编号。例如，如果总体具有 500 个元素，我们按顺序将它们以三位数进行编号，从 001 开始到 500 结束。假设我们想要一个大小为 50 的简单随机样本。在这种情况下，利用计算机随机数生成器或者随机数表，我们可以产生一系列三位的随机数。然后我们将这些随机数与总体中元素的编号相匹配（将相同编号的元素取出），直到我们选出大小为 50 的样本为止。

有时，我们无法为总体中所有的元素编号（或识别出所有的元素）。在这种情况下，我们经常采用系统抽样（systematic sampling）的方法。通过系统抽样，我们每隔 k 个元素进行抽样，直至我们选出所需大小的样本。由这种方法所选出的样本近似是随机的。实际抽样的情况可能需要我们采用近似随机的样本。

假设电信设备的分析师对于一个电信设备客户的随机样本进行调查来确定平均设备支出情况。样本均值将提供该分析师一个对于总体支出均值的估计。样本均值和总体均值的任何差别都

被称为抽样误差 (sampling error)。

- **抽样误差的定义。**抽样误差是统计量的观测值和所要估计的真实值之间的差别。

一个随机样本以无偏的方式反映了总体的性质，而样本统计量，如样本均值，是基于随机样本计算得到的，它们是对应总体参数的有效估计量。

一个样本统计量是一个随机变量。换句话说，不仅总体中的原始数据具有一个分布，而且样本统计量也具有一个分布。这个分布就是统计量的抽样分布。

- **一个统计量抽样分布的定义。**一个统计量的抽样分布是指从相同总体中以相同数目随机抽取出的样本计算得到的统计量所能够取到的所有不同可能值的分布。

例如，在样本均值的例子中，我们将其称为“样本均值的抽样分布”或者样本均值的分布。我们将在本章的后面多次提到抽样分布。然而，下面我们要了解一下另一个在投资分析中有用的抽样方法。

6.2.2 分层随机抽样

刚才讨论的简单随机抽样方法可能并不是在所有情况下都是最好的方法。我们经常使用的另外一种方法是分层随机抽样。

- **分层随机抽样 (stratified random sampling) 的定义。**在分层随机抽样中，总体会根据一个或者多个分类标准被分为许多子总体 (层)。然后我们根据每个层在总体中所占的相对大小比例，从中抽取简单随机样本。最后将所有这些从每一层中抽取的样本放在一起就形成了一个分层随机样本。

不同于简单随机抽样，分层随机抽样保证了所关心的总体子类都在样本中得到体现。而另一个优点是由分层随机抽样所产生的参数估计值相对于由简单随机抽样所得到的估计值具有更高的精度，即较小的方差和离散度。

债券指数化是分层抽样的一个常见的应用领域。指数化 (indexing) 是一种投资策略，在这种策略中，投资者建立一个投资组合以反映一个特定指数的表现。在纯债券指数化中，同样也被称为完全复制方法，投资者试图通过以它们所占的市场价值权重的比例持有指数中的所有债券来完全复制指数。然而，许多债券指数包含成千上万的发行债券，因此纯债券指数化是很难实行的。另外，因为许多债券没有流动性很高的市场，所以这样做的交易成本会很高。虽然一个简单随机样本可能是一个解决成本问题的方法，但是样本可能会与指数中主要的风险因子 (例如利率敏感性) 不匹配。因为固定收益投资组合中主要的风险因子是为人们所熟知的，而且可以定量化，因此，分层抽样提供了一个更加有效的方法。在这种方法中，我们将指数债券总体分成相同久期 (利率敏感性)、现金流分布、行业、信用水平和敞口情况的一些组类。我们就将每个组类作为一个层或者一个单元 (在这一环境中经常使用的一个术语)。[⊖]于是，我们从每个层中选择被复制指数中该层所占市场权重数量的样本。

⊖ 参见 Fabozzi (2004b)。

例 6-1 债券指数和分层抽样

假设你是一家跟踪雷曼兄弟政府债券指数的共同基金的经理。你现在正在采用不同的方法来进行指数化,其中包括一种分层抽样的方法。你首先从美国国库券中区分出机构债券。对于这两组中的每一个,你定义 10 个到期日的区间——1 到 2 年,2 到 3 年,3 到 4 年,4 到 6 年,6 到 8 年,8 到 10 年,10 到 12 年,12 到 15 年,15 到 20 年以及 20 到 30 年——并且你也将(年利率)小于等于 6% 的付息债券与大于 6% 的付息债券相分离。

(1) 这个抽样计划需要多少单元或者层数?

(2) 如果你使用这个抽样计划,那么指数化的投资组合中所含有的最小的发行债券数量是多少?

(3) 假设在每个单元中满足选择要求的证券中选择时,你使用一个关于证券市场流动性的标准。这时所得到的样本是随机的吗?解释你的结论。

解:

(1) 我们有 2 个债券发行者的分类、10 个到期日的分类和 2 个息票利率的分类。因此,该计划共有 $2(10)(2) = 40$ 个不同的层数或单元。(这个回答应用了概率论中的一些概念一章中所讨论的乘法计数法则。)

(2) 你不能在每个单元中具有少于 1 个的发行债券,所以这个投资组合必须至少包含 40 个发行债券。

(3) 如果你对于单元中证券的选择应用任何其他标准,那么每一个单元中所包含的证券将不会以相同的概率被选取。结果,这种抽样不是随机的。在实际应用中,应用分层抽样的指数化通常并不严格产生随机样本,因为每个单元中发行债券的选择受到不同其他标准的影响。因为在这个应用中抽样的目的不是对于总体参数进行推断,而是指数化一个投资组合,在这个分层抽样的应用中缺乏随机性本身并不是一个问题。

在下面一节中,我们将讨论金融分析师在抽样和选择样本过程中所产生的实际问题中所遇到的不同数据类型。

6.2.3 时间序列数据和横截面数据

投资分析师常常与时间序列和横截面数据打交道。一个时间序列是一组在离散的等间距时点上收集的收益率序列(例如一组股票月收益率的历史序列)。横截面数据是关于个体、组类、地理区域或者公司在某一时点的特征的数据。2003 年年末所有纽约证券交易所上市公司的每股账面价值是横截面数据的一个例子。

经济或者金融理论并没有提供我们选择样本时间区间长短的理论基础。作为分析师,我们可能不得不去寻求细微的判断线索。例如,将一个固定汇率区间段的数据和浮动率区间段的数据相合并是不恰当的。当汇率固定时,汇率的方差肯定小于当汇率允许浮动时的情况。因此,我们不会从一个以单一参数描述的总体中进行抽样。^①紧缩的和宽松的货币政策同样会对于股票的收益率分布产生影响;因此将紧缩货币政策和宽松货币政策区间的数据合并也是不合适的。例 6-2 举例说明了当从多个分布中抽样时所产生的问题。

① 当一个时间序列的均值和方差随时间变化时,时间序列不是平稳的。我们在时间序列分析一章中会更加具体地讨论平稳性。

例 6-2 计算夏普比率：使用一年还是两年的季度数据？

分析师经常使用夏普比率来评估所管理的投资组合的业绩表现。夏普比率（Sharpe ratio）是超出无风险收益率的平均收益率除以收益率的标准差的值。这个比率度量了每单位收益率标准差所获得的超额收益率。

为了计算夏普比率，假设一位分析师收集了 8 个季度的超额收益率（即超出无风险收益率的总收益率）。在第 1 年中，投资组合的投资经理遵循了一个低风险的策略，而在第 2 年中，该经理采取了一个高风险的策略。对于这两年的每一年，分析师都对于超出某个用来评价经理基准的超额季度收益率进行了跟踪。

对于这两年的每一年，基准的夏普比率是 0.21。表 6-1 给出了投资组合夏普比率的计算。

对于第 1 年，在这一年中经理遵循了一个低风险策略，超过无风险收益率的平均季度收益率为 1%，标准差为 4.62%。因此，夏普比率为 $1\% / 4.62\% = 0.22$ 。第 2 年的结果与第 1 年

表 6-1 夏普比率的计算：低风险和高风险的策略

季度/度量	第 1 年的超额收益率	第 2 年的超额收益率
第一季度	-3%	-12%
第二季度	5	20
第三季度	-3	-12
第四季度	5	20
季度平均值	1%	4%
季度标准差	4.62%	18.48%
夏普比率 = $0.22 = 1 / 4.62 = 4 / 18.48$		

类似，除了它具有更高的平均收益率和波动率。第 2 年的夏普比率为 $4 / 18.48 = 0.22$ 。在第 1 年和第 2 年中基准的夏普比率是 0.21。因为较大的夏普比率要比较小的夏普比率更好（给予每单位风险更多的收益率），所以该经理表现出超过基准的业绩表现。

现在，假设分析师认为较大的样本要比较小的样本要好。因此，她决定将两年的数据放在一起，并计算 8 个季度观测值的夏普比率。这两年的季度平均超额收益率为每年平均超额收益率的平均值。对于这两年的区间，平均超额收益率为每季度 $(1 + 4) / 2 = 2.5\%$ 。由样本均值 2.5% 所求得的所有 8 个季度的标准差为 12.57%。该投资组合在这两年时间区间中的夏普比率现在为 $2.5 / 12.57 = 0.199$ ；这时基准的夏普比率仍是 0.21。因此，当将这两年的收益率放在一起时，该经理似乎比基准以及之前两年分开计算的结果提供了每单位风险更少的收益率。

利用 8 个季度收益率数据的问题在于分析师违背了抽样收益率必须来自于同一总体的假设。由于该经理的投资策略发生了改变，在第 2 年的收益率遵循一个与第 1 年收益率不同的概率分布。显然，在第 1 年，收益率是由小于第 2 年均值和方差的总体产生的。将第 1 年和第 2 年的结果合并在一起所产生的样本不能代表任何一个总体。因为更大的样本不满足模型的假设，分析师由这个更大的样本所得到的任何结论都是不正确的。对于这个例子，相对于较大的样本，她利用较小的样本会得到更好的结果，因为较小的样本反映了相同的收益率分布。

第 2 类基本的数据类型是横截面数据。^①在横截面数据中，样本的观测值反映个体、组类、地理区域或者公司在某个时点上的某个特征。先前讨论的电信行业分析师实际上收集的就是下一年

① 读者也可能遇到两种类型都存在的数据集，即既有时间序列数据的一面，也有横截面数据的一面。面板数据（panel data）包含了多个观测单位在不同时点的单个特征的观测值。例如，欧元区国家在一个 5 年时间段上的年通货膨胀率数据将以面板数据来表现。纵向数据（Longitudinal data）包含相同观测单位在不同时点上的特征的观测值。单个公司在 10 年的时间段上的一组金融比率观测值是纵向数据的一个例子。面板数据和纵向数据也可以用数组（矩阵）来表示，其中后面的行表示在下一时间区间上的观测值。

计划资本支出的横截面数据。

无论我们在哪个时点上进行横截面抽样，只要我们希望将那个数据以有意义的方式进行汇总，特定的假设是一定会满足的。另外，一个有用的方法是将所关心的观测值认为是一个来自于某个给定均值和方差的对应总体的随机变量。在我们收集样本并开始汇总数据时，我们必须确保所有的数据事实上确实是来自于所对应的相同总体。例如，一位分析师可能对于公司应用它们存货资产的效率情况感兴趣。然而，因为不同的经营环境，一些公司比另一些公司具有更快的存货周转率（例如，通常食品杂货店的存货周转速度要快于汽车制造商），所以存货周转率的分布可能不是以一个给定均值和方差的单一分布所描述的。因此，汇总所有公司的存货周转率可能是不恰当的。如果随机变量是有不同的对应分布所产生的，那么由这个合并样本所计算的样本统计量是和某个对应的总体参数无关的。在这种情况下，抽样错误的大小是不能获知的。

诸如这些例子中，分析师经常按行业来汇总公司水平上的数据。试图汇总按行业汇总部分数据回答了不同对应分布的问题，但是大公司可能在不只一个行业部门中经营，所以分析师应该确保他们知道如何将公司分配到每个行业组别中。

无论我们处理的是时间序列数据还是横截面数据，我们必须确保使用的是代表我们希望研究总体的随机样本。有了代表性样本中所推断信息的目标，我们现在转向本章的下面一部分，将关注点聚焦于中心极限定理和总体均值的点估计和区间估计。

6.3 样本均值的分布

在本章的前面，我们介绍了一个电信设备分析师的例子，该分析师决定通过抽样来估计用户计划资本支出的均值。假设样本代表了所对应的总体，这位分析师如何来评估其估计总体均值的抽样误差呢？样本均值本身是一个具有概率分布的随机变量，它被视为一个随机变量随机结果函数的一个公式。这个概率分布被称为统计量的抽样分布。[⊖]为了估计预期的样本均值与对应的总体均值到底有多接近，分析师需要理解均值的抽样分布。幸运的是，我们有一个结果（中心极限定理），它能够帮助我们理解我们所面对的许多估计问题的均值抽样分布。

中心极限定理

作为概率论中在实践上最有用的定理之一，中心极限定理对于我们如何构建置信区间和检验假设具有重要的意义。正式地，该定理被表述如下：

- 中心极限定理（central limit theorem）。给定一个具有均值 μ 和有限方差 σ^2 的任一概率分布的总体，当样本量 n 很大时，由总体中抽取的样本大小为 n 的样本均值 $\bar{X}$ 的抽样分布将近似服从均值为 μ （总体均值）、方差为 σ^2/n （总体方差除以 n ）的正态分布。

中心极限定理使得我们可以利用样本均值来精确地做出对于总体均值的概率估计，而无论总体的分布是什么，因为样本均值在大样本的情况下服从近似正态分布。一个显然的问题是，“什么时候样本大小才算足够大以使我们假设样本均值服从正态分布呢？”一般来说，当样本量

⊖ 因为“样本均值”也在其他的意义下被使用，有时这会引入概念混淆。当我们对于一个特定样本计算其样本均值时，我们获得一个确定的数字，比如说8，如果我们认为“样本均值是8”，我们使用“样本均值”表示作为一个随机变量的样本均值的一个特定结果。数字8当然是常数而且它不具有概率分布。在这个讨论中，我们并没有将“样本均值”理解为与一个特定样本相关的常数。

n 大于等于 30 时, 我们可以假设样本均值近似服从正态分布。[⊖]

中心极限定理认为样本均值分布的方差为 σ^2/C 。方差的正的平方根是标准差。一个样本统计量的标准差被认为是统计量的标准误。样本均值的标准误是一个在实践中应用中心极限定理时所用的重要数值。

- 样本均值标准误 (standard error of the sample mean) 的定义。对于由一个标准差为 σ 的总体中所取出样本所计算的样本均值 $\bar{X}$, 其标准误可以用下面两个表达式之一来表示:

当我们已知 σ (总体标准差) 时,

$$\sigma \bar{X} = \frac{\sigma}{\sqrt{n}} \quad (6-1)$$

或者当我们不知道总体标准差并需要利用样本标准差 s 来估计它时,[⊖]

$$s \bar{X} = \frac{s}{\sqrt{n}} \quad (6-2)$$

在实际应用中, 我们几乎总是需要使用式 (6-2)。s 的估计值是由样本方差 s^2 的平方根给出的, 具体计算如下:

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1} \quad (6-3)$$

我们将马上看到如何通过使用置信区间的技术来利用样本均值和其标准误对总体均值进行概率估计。然而, 首先我们会提供一个例子来说明中心极限定理的威力。

例 6-3 中心极限定理

我们都知道了无论对应的总体分布是什么, 大样本的样本均值都将会服从正态分布。为了介绍中心极限定理的运作方式, 我们在这个例子中设定一个明显的非正态分布, 并利用它来产生大量大小为 100 的随机样本。然后我们计算每个样本的样本均值。所计算的样本均值的频数分布是在这个样本大小下样本均值抽样分布的一个近似。这个抽样分布与一个正态分布看上去是否很相像?

我们回到研究电信行业业务资本支出计划的分析师的例子当中。假设电信设备的资本支出来自一个下界为 0 美元、上界为 100 美元的连续均匀随机变量, 简单地称其为一个 (0, 100) 的均匀随机变量。这个连续均匀随机变量的概率分布具有一个相当简单的形状 (不是正态分布)。它是一条水平直线, 垂直截距等于 1/100。不同于一个正态随机变量 (其结果接近于均值的可能性很大), 对于一个均匀随机变量来说所有可能结果都是等可能的。

⊖ 当对应的总体偏离正态分布很远时, 样本大小可能要求超过 30, 以使得正态分布能够成为均值抽样分布的一个很好的描述。

⊖ 我们需要注意的一个技术点是: 当我们从大小为 N 的有限总体中选取大小为 n 的样本时, 我们使用了一个缩减因子来估计样本均值标准误的估计值, 它被称为有限总体修正因子 (fpc)。fpc 等于 $[(N - n)/(N - 1)]^{1/2}$ 。因此, 如果 $N = 100$ 而 $n = 20$, $[(100 - 20)/(100 - 1)]^{1/2} = 0.898\ 933$ 。如果我们已经根据式 (6-1) 或者式 (6-2) 估计出了一个标准误, 比如说为 20, 那么新的估计值为 $20(0.898\ 933) = 17.978\ 663$ 。fpc 只有当我们从有限总体中无放回地抽样时才会应用; 当样本大小 n 相对于 N 很小 (例如, 小于 N 的 5%), 大部分从业者也不会应用 fpc。关于有限总体修正因子的更多内容参见 Daniel 和 Terrell (1995)。

为了展现中心极限定理的威力，我们采用一个蒙特卡罗模拟来研究电信业务的资本支出情况。^①在这个模拟中，我们收集了200个100家公司的资本支出的随机样本（200个随机抽取，每个由100（ $n=100$ ）家公司的资本支出组成）。在每次模拟试验中，资本支出的100个值被从（0，100）的均匀分布中产生。对于每个随机样本，我们可以计算其样本均值。我们总共进行了200次模拟试验，因为我们已经设定了产生样本的分布，我们已知资本支出总体的均值等于（\$0 + \$100）/2 = \$50（百万）；资本支出总体的方差等于 $(100 - 0)^2 / 12 = 833.33$ ；因此，标准差是28.87美元，而在中心极限定理下的标准误是 $28.87 / \sqrt{100} = 2.887$ 。^②

这一蒙特卡罗试验的结果以频数分布的形式列示在表6-2中了。这个分布是样本均值估计的抽样分布。

表6-2 频数分布：200个（0，100）均匀随机变量所产生的随机样本

样本均值的范围（1 000 000 美元）	绝对频数	样本均值的范围（1 000 000 美元）	绝对频数
$42.5 \leq \bar{X} < 44$	1	$50 \leq \bar{X} < 51.5$	39
$44 \leq \bar{X} < 45.5$	6	$51.5 \leq \bar{X} < 53$	23
$45.5 \leq \bar{X} < 47$	22	$53 \leq \bar{X} < 54.5$	12
$47 \leq \bar{X} < 48.5$	39	$54.5 \leq \bar{X} < 56$	12
$48.5 \leq \bar{X} < 50$	41	$56 \leq \bar{X} < 57.5$	5

注： $\bar{X}$ 是每个样本平均的资本支出。

频数分布可以被描述成钟型的，并且中心接近于总体均值50。出现最频繁的或者众数集中的范围为48.5~50的区间，共有41个观测值。样本均值的总体平均值为49.92美元，其标准误为2.80美元。所计算出的标准误接近于中心极限定理所给出的2.887美元。计算出的和在中心极限定理中所期望的均值和标准差之间的差别就是随机性（样本误差）的结果。

总之，虽然对应总体的分布偏离正态分布很远，模拟的结果表明一个正态分布能够很好地描述样本均值的估计抽样分布，其均值和标准误都和由中心极限定理所预测的值相一致。

总结一下，根据中心极限定理，当我们从任一分布中取样时，只要我们的样本容量很大，样本均值的分布将会具有下列的性质：

- 样本均值 $\bar{X}$ 的分布将是近似正态的。
- $\bar{X}$ 分布的均值将等于抽取样本的总体的均值。
- $\bar{X}$ 分布的方差将等于总体方差除以样本大小。

有了中心极限定理，我们下面要讨论涉及总体参数估计的概念和工具，其中特别关注总体均值的估计。我们关注总体，是因为分析师更经常遇到的是总体均值的区间估计，而不是其他类型的区间估计。

① 蒙特卡罗模拟涉及利用一个计算机来反映面临风险的一个系统的运作。蒙特卡罗模拟的重要部分是从一个特定的概率分布或一组概率分布中产生大量的随机样本。

② 如果 a 是均匀随机变量的下界， b 是其上界，那么该随机变量的均值由 $(a + b) / 2$ 给出，其方差由 $(b - a)^2 / 12$ 给出。在常见概率分布一章中对于连续均匀随机变量已经进行了充分的讨论。

6.4 总体均值的点估计和区间估计

传统上,统计推断包含两个分支——假设检验和参数估计。假设检验回答“这个参数的值(比如一个总体均值)是否等于某个特定值(比如0)?”这个问题。在这个过程中,我们有一个关于参数值的假设,并且我们试图确定由样本获得的证据支持或者不支持那个假设。我们将在假设检验一章中详细讨论假设检验。

统计推断的第2个分支就是参数估计,并且它是本章所关注的重点。参数估计试图回答“参数值(例如总体均值)到底是多少?”的问题。在参数估计中,不同于假设检验,我们不从一个关于参数值的假设开始着手并检验该假设。相反,我们试图最大限度地利用样本中的信息来形成不同类型中的参数估计值中的一个。在参数估计中,我们希望找到一个最好的计算出的单个数字来对于未知的总体参数进行估计(一个点估计)。与计算点估计类似,我们也可能希望计算一个在一定概率水平下包含未知总体参数的范围值(一个置信区间)。在4.1节中,我们讨论参数的点估计,而在4.2节中,我们介绍总体均值置信区间的构造。

6.4.1 点估计量

本章中所介绍的一个重要概念是样本统计量是一个随机变量,它可以看做与随机结果有关的公式。我们用来计算样本均值和所有其他样本统计量的公式是估计公式或者是估计式(estimators)。我们利用估计量从样本观测值中所计算出的特定值被称为一个估计值(estimate)。一个统计量具有一个抽样分布;一个估计值是一个关于给定样本的固定数字,因此,它不具有抽样分布。举一个均值的例子,根据给定样本所计算的样本均值,它是用来作为总体均值的估计值,这个估计值被称为总体均值的点估计(point estimate)。正如例6-3所显示的,由于不同的样本是从总体中抽取出来的,样本均值的公式能够得到重复抽样的样本的不同结果。

在许多应用中,我们要在许多给定参数的可能估计量中做出一个选择。那么我们如何做出我们的选择呢?我们经常对于估计量进行选择,因为我们希望它们具有一个或者多个良好的统计性质。下面就是对于估计量3个良好性质的简要介绍:无偏性(没有偏差的),有效性和相合性。^①

- 无偏性(unbiasedness)的定义。一个无偏估计量是指其期望值(它抽样分布的均值)等于它所要估计的参数值。

例如,样本均值 $\bar{X}$ 的期望值等于总体均值 μ ,所以我们认为样本均值是(总体均值的)一个无偏估计量。样本方差 s^2 是利用除数 $n-1$ 计算得到的(式(6-3)),它是总体方差 σ^2 的一个无偏估计量。如果我们利用除数 n 来计算样本方差,那么该估计量将会是有偏的:它的期望值将小于总体方差。我们将由除数 n 计算得到的样本方差称为总体方差的一个有偏估计量。

当我们可以找到参数的一个无偏估计量时,我们通常就可以找到许多其他的无偏估计量。那么我们如何在这些无偏估计量中进行选择呢?有效性的标准提供了我们一种从许多参数的无偏估计量中进行选择的方法。

- 有效性(efficiency)的定义。当对于同一个参数,没有其他无偏估计量具有比该估计量更小的抽样分布方差时,我们就称该无偏估计量是有效的。

^① 对于这些估计量性质的更加详细的处理,参见 Daniel 和 Terrell (1995) 或者 Greene (2003)。

为了解释这个定义，在重复抽样的样本中，我们认为一个有效估计量的估计值应该比其他无偏估计量更加紧密地聚集在均值的周围。有效性是估计量的一个重要性质。^①样本均值 $\bar{X}$ 是总体均值的有效估计量。而样本方差 s^2 是总体方差 σ^2 的有效估计量。

回忆一下，一个统计量的抽样分布是根据给定的样本容量定义的。不同的样本容量定义了不同的抽样分布。例如，样本均值抽样分布的方差对于大样本来说较小。无偏性和有效性是对于任何大小样本估计量抽样分布的普遍性质。一个估计量在样本大小为 10 时是无偏估计量，在样本大小为 1 000 时，也同样是无偏估计量。然而在一些问题中，我们不能找到具有这种良好性质的估计量，如在小样本中的无偏性。^②在这种情况下，统计学家可以对于在极大样本中基于估计量抽样分布的这些性质所选出的估计量进行调整，这就是所谓的估计量的渐近性质。在这些性质中，最重要的是相合性。

- 相合性 (consistency) 的定义。一个相合估计量是指随着样本量的增大，该估计量接近总体参数真值的概率也会增大，并趋近于 1。

稍稍技术化一点，我们可以定义一个相合估计量等于这样的估计量，其抽样分布随着样本容量趋近于无穷，会集中到它要估计参数的真值之上。样本均值，除了是一个有效估计量之外，同样也是总体均值的一个相合估计量：随着样本容量 n 趋向于无穷大，其标准误 $\sigma/\sqrt{n}$ 趋向于 0，所以其抽样分布正好集中于总体均值 μ 一点之上。总结一下，我们可以将一个相合估计量认为是一个随着我们样本容量的扩大，而趋向于产生关于总体参数越来越精确估计值的估计量。如果一个估计量是相合的 (consistent)，我们可以试图利用更大的样本来对于估计值进行计算，从而增加总体参数估计值的准确性。然而，对于一个不具有相合性的估计量来说，增加样本容量并不能对获得准确估计值概率的提高带来帮助。

6.4.2 总体均值的置信区间

当我们需要一个数值作为总体参数的一个估计量时，我们利用的是点估计。然而，因为存在抽样误差，点估计是不可能完全等于任一给定样本所代表总体的参数值的。通常，一个比点估计更加有用的方法是去寻找一个在设定概率水平之下包含该参数期望值的区间范围——参数的一个区间估计。置信区间就扮演了这样的角色。

- 置信区间 (confidence interval) 的定义。一个置信区间是一个区间范围，我们所要估计的参数将以一个给定的概率 $1 - \alpha$ (其被称为置信度 (degree of confidence)) 被包含在这个区间之中。

置信区间边界的端点被称为置信区间的下界和上界。在本章中，我们只讨论双边的置信区间，我们会同时计算下界和上界的置信区间。^③

置信区间经常被赋予概率的或者实际的解释。在概率解释中，我们将一个总体均值 95% 的置信区间解释如下：在反复抽样过程中，从长期来看，有 95% 的这种置信区间将包含或包括总体均值。例如，假设我们总体中抽样 1 000 次，并基于每个样本利用所计算的样本均值建

① 一个有效估计量有时也称为最佳无偏估计量。

② 这个问题通常发生在回归和时间序列分析中，我们将在后面的章节中对此进行讨论。

③ 我们也可以定义出对于总体参数的两种单边置信区间。一个下侧单边置信区间只具有一个下界。与这个置信区间相关的论断是总体参数会以给定的置信水平大于或等于下界。一个上侧单边置信区间只有一个上界；相关的论断是总体参数小于或者等于该上界的概率等于所设定的置信度。然而，投资研究者们很少用单边置信区间进行表达。

立一个 95% 的置信区间。由于随机性，这些置信区间将各不相同，但是我们期望 95% 或者其中的 950 个区间将包含未知的总体均值的真值。在实际应用中，我们一般不进行这样的反复抽样。因此，在实际解释中，我们认为对于一个 95% 的置信区间，我们具有 95% 的信心，相信它包含了总体的均值。我们有充分的理由做出这一表述，因为我们知道所有由这一相同方法所构造的可能的置信区间中有 95% 将包含总体均值。在本章中所讨论的置信区间具有和下面这个基本形式类似的结构：

- **置信区间的构建。** 对于一个参数的 $(1 - \alpha)\%$ 置信区间具有下面的结构：

$$\text{点估计} \pm \text{可靠性因子} \times \text{标准误}$$

式中

点估计 = 参数的一个点估计（样本统计量的一个值）

可靠性因子 = 基于点估计的假设分布和置信区间中置信度 $(1 - \alpha)$ 的一个数

标准误 = 提供点估计的样本统计量的标准误[⊖]

当我们从一个已知方差的正态分布中进行抽样时，就产生了最基本的对于总体均值的置信区间。在这种情况下可靠性因子是基于均值为 0、方差为 1 的标准正态分布得到的。一个标准正态随机变量一般记为 Z 。记号 z_{α} 指的是标准正态分布中使得右侧尾部概率为 α 的那一点的值。例如，标准正态随机变量的取值有 0.05 或者 5% 的概率要比 $z_{0.05} = 1.65$ 要大。

假如我们想要构造一个总体均值的 95% 置信区间，为了这个目的，我们已经从一个已知方差为 $\sigma^2 = 400$ （因此， $\sigma = 20$ ）的正态分布中抽取了一个大小为 100 的样本。我们计算得到样本均值 $\bar{X} = 25$ 。因此，我们对于总体均值的点估计等于 25。如果我们从正态分布的均值向上移动 1.96 个标准差，那么在右侧尾部的概率将为 0.025 或者 2.5%。由正态分布的对称性，如果我们从均值向下移动 1.96 个标准差，那么在左侧尾部的概率也将变为 0.025 或者 2.5%，故总共有 0.05 或者 5% 的概率在两个尾部，而 0.95 或者 95% 的概率位于中间的部分。因此， $z_{0.025} = 1.96$ 是这个 95% 置信区间的可靠性因子。注意对于置信区间的 $(1 - \alpha)\%$ 和可靠性因子 $z_{\alpha/2}$ 之间的关系。由式 (6-1) 所给出的样本均值的标准误为 $\sigma_{\bar{X}} = \frac{20}{\sqrt{100}} = 2$ 。因此，置信区间具有下界 $\bar{X} - 1.96\sigma_{\bar{X}} = 25 - 1.96(2) = 25 - 3.92 = 21.08$ 。置信区间的上界等于 $\bar{X} + 1.96\sigma_{\bar{X}} = 25 + 1.96(2) = 25 + 3.92 = 28.92$ 。对于总体均值 95% 的置信区间为 21.08 ~ 28.92。

- 总体均值的置信区间（confidence intervals for the population mean）（方差已知的正态分布总体）（normally distributed population with known variance）。当我们从一个已知方差 σ^2 的正态分布总体中抽样时，总体均值 μ 的一个 $(1 - \alpha)\%$ 置信区间为：

$$\bar{X} \pm z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \tag{6-4}$$

下面给出了最常用的置信区间可靠性因子。

- 基于标准正态分布的置信区间可靠性因子（reliability factors for confidence intervals based on the standard normal distribution）。当我们基于标准正态分布建立置信区间时，我们利用下列可靠性因子：[⊖]

⊖ （可靠性因子） $\times$ （标准误）这个量有时也被称为估计量的精度；这个乘积值越大就意味着总体参数估计的精度越低。

⊖ 大部分从业者使用保留到小数点后两位的 $z_{0.05}$ 和 $z_{0.005}$ 。作为参考，更加精确的 $z_{0.05}$ 和 $z_{0.005}$ 分别为 1.645 和 2.575。如果要快速计算 95% 置信区间， $z_{0.025}$ 有时被从 1.96 近似为 2。

- 90% 置信区间：使用 $z_{0.05} = 1.65$
- 95% 置信区间：使用 $z_{0.025} = 1.96$
- 99% 置信区间：使用 $z_{0.005} = 2.58$

这些可靠性因子显示了所有置信区间的一个重要事实。随着我们增加置信度时，置信区间将会变宽，这使得我们无法获得我们所要估计的量的精确信息。因此，“想要确信程度越高，那么只能确信的信息越少。”[⊖]

在实际应用中，样本均值的抽样分布至少服从近似正态的假设通常是合理的，这或者是因为所关心的分布本身是近似正态的，或者是因为我们具有大样本并应用了中心极限定理。然而，我们很少在实际应用中知道总体的方差。当总体方差是未知的，但样本均值至少服从近似正态分布时，我们有两种可接受的计算总体参数置信区间的处理方法。我们将马上讨论其中一种更加保守的方法，它是基于学生氏 t 分布（简称为 t 分布）。[⊗]在投资学文献中，它是在总体方差未知时对于均值估计和假设检验的最常用方法，而无论样本容量是大还是小。

第2种构造总体均值置信区间的方法为基于标准正态分布的 z 替代方法。它只有当样本量很大时才被使用。（一般，样本容量大于等于 30 才被视为大样本。）与式（6-4）给出的置信区间不同，这个置信区间在计算样本均值标准误时使用的是样本标准差 s （式（6-2））。

- 总体均值的置信区间—— z 替代方法（confidence intervals for the population mean—the z -alternative）（大样本，样本方差未知（large sample, population variance unknown））。当我们从一个方差未知的正态分布总体中抽样且所抽取的样本是大样本时，总体均值 μ 的一个 $(1 - \alpha)\%$ 置信区间为：

$$\bar{X} \pm z_{\alpha/2} \frac{s}{\sqrt{n}} \quad (6-5)$$

因为这类置信区间经常出现，我们在例 6-4 中介绍一下它的计算。

例 6-4 夏普比率总体均值的置信区间—— z -统计量

假设一位投资分析师从美国股权共同基金中选取了一个随机样本，并计算出了平均的夏普比率。该样本容量为 100，并且平均的夏普比率为 0.45。该样本具有的标准差为 0.30。利用一个基于标准正态分布的可靠性因子，计算并解释所有美国股权共同基金总体均值的 90% 置信区间。

正如之前所给出的，这个 90% 的置信区间的可靠性因子为 $z_{0.05} = 1.65$ ，故置信区间为

$$\bar{X} \pm z_{0.05} \frac{s}{\sqrt{n}} = 0.45 \pm 1.65 \frac{0.30}{\sqrt{100}} = 0.45 \pm 1.65(0.03) = 0.45 \pm 0.0495$$

置信区间为 0.4005 ~ 0.4995，或者保留到小数点后两位的 0.40 ~ 0.50。分析师可以说有 90% 的信心认为这个区间包含了总体均值。

在这个例子中，分析师没有对于描述总体的概率分布作任何具体的假设。相反，分析师依靠中心极限定理产生了一个样本均值的渐近正态分布。

⊖ Freund 和 Williams (1997), p. 266。

⊗ 统计量 t 的分布被称为学生氏 t 分布，它是为了纪念 W. S. Gosset（他在 1908 年发表了关于这个分布的研究结果）所使用的笔名“学生”（student）而来的。

例6-4表明,即使我们对于所关心的总体分布情况不确定,我们仍旧可以构造总体均值的置信区间,只要样本容量足够大,因为在这种情况下我们可以应用中心极限定理。

我们现在要转向保守(传统)的一种方法,当总体方差未知时,使用 t 分布来构造总体均值的置信区间。对于基于样本服从方差未知的正态分布总体的置信区间来说,理论上正确的可靠性因子是基于 t 分布的。利用基于 t 分布获得的可靠性因子对于一个小样本来说是非常重要的。当总体方差未知时,即使我们所拥有的是一个样本而且可以使用经过中心极限定理证明的 z 可靠性因子,使用一个 t 可靠性因子也是恰当的。在大样本的例子中, t 分布提供了一个更加保守的(更宽的)置信区间。

t 分布是一个由单一参数(称为自由度(degrees of freedom, df))所定义的对称概率分布。对于每个自由度的数值都定义了这类分布族中的一个分布。我们不久将对 t 分布和标准正态分布进行一个比较,但是首先我们需要理解自由度的概念。我们可以通过对于样本方差计算的检验来理解这个概念。

式(6-3)给出了我们所使用的样本方差的无偏估计量。在分母上的项 $n-1$ 是样本容量减去1,它是利用式(6-3)估计总体方差时所用的自由度数值。我们也使用 $n-1$ 作为确定基于 t 分布的可靠性因子的自由度数值。我们使用“自由度”这一术语是因为在一个随机样本中,我们假设观测值的选择之间是独立的。而样本方差的分子使用了样本均值。那么样本均值的使用会如何影响样本方差计算公式中观测值独立选择数值的呢?例如,在样本容量为10、样本均值为10%的情况下,我们只可以自由地选择9个观测值。无论这9个观测值如何选取,我们总能求解出第10个观测值,使这10个观测值的均值等于10%。而这从样本方差公式的角度来看,就是具有9个自由度。给定我们必须首先从总共 n 个独立观测值中计算出样本均值,那么在计算样本方差时,只有 $n-1$ 个观测值可以被独立地选择。自由度的概念在统计学中经常出现,你将会在后面的章节中经常遇到它。

假设我们从一个正态分布中抽样。比率 $z = (\bar{X} - \mu) / (\sigma / \sqrt{n})$ 是服从均值为0、标准差为1的正态分布;然而,比率 $z = (\bar{X} - \mu) / (s / \sqrt{n})$ 服从均值为0,自由度为 $n-1$ 的 t 分布。有 t 表示的比率并不服从正态分布,因为 t 是两个随机变量样本均值和样本标准差的比率。标准正态随机变量的定义只涉及一个随机变量,即样本均值。

然而,随着自由度的增加, t 分布趋近于标准正态分布。图6-1给出了标准正态分布和两个 t 分布,一个具有 $df=2$,而另一个 $df=8$ 。

图6-1所显示的3个分布中,标准正态分布明显峰度最高;它的尾部趋近于0的速度要快于两个 t 分布。和正态分布一样, t 分布也是关于均值0对称的分布。在这3个分布中,具有 $df=2$ 的 t 分布的峰度最低,它的尾部要高于正态分布和 $df=8$ 的 t 分布。 $df=8$ 的 t 分布具有中等程度的峰度,它的尾部高于正态分布,但低于 $df=2$ 的 t 分布。随着自由度的增加, t 分布趋近于标准正态分布。 $df=8$ 的 t 分布比 $df=2$ 的 t 分布更接近于标准正态分布。

在均值正负4个标准差的区间之外,标准正态分布曲线下的面积近似趋近于0;然而两个 t 分布在4个标准差之外的每条曲线下仍旧存在一些面积。 t 分布具有更厚的尾巴,但是 $df=8$ 的 t 分布的尾巴要更接近于正态分布。随着自由度的增加, t 分布尾部的厚度会逐渐变小。

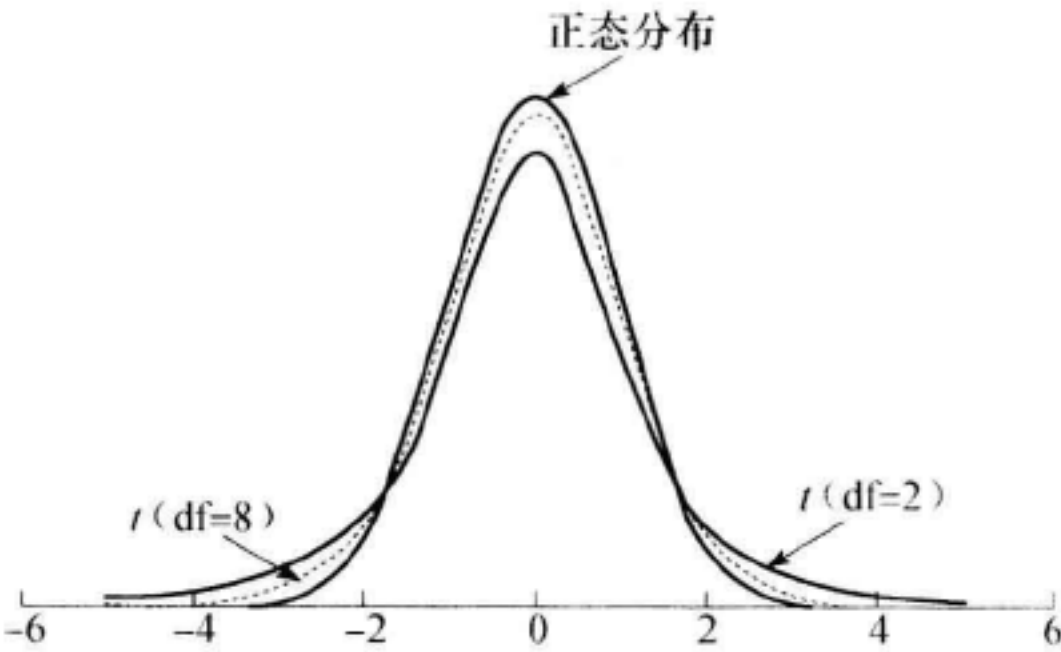


图6-1 学生氏 t 分布与标准正态分布

t 分布中经常使用的值我们都在书后的表格中予以列出。对于每个自由度来说, 5 个值是给定的: $t_{0.10}$, $t_{0.05}$, $t_{0.025}$, $t_{0.01}$ 和 $t_{0.005}$ 。 $t_{0.10}$, $t_{0.05}$, $t_{0.025}$, $t_{0.01}$ 和 $t_{0.005}$ 的值分别表示在给定具体的自由度数时, 右侧尾部的概率为 0.10, 0.05, 0.025, 0.01 和 0.005 的值。[⊖] 例如, 对于 $df = 30$, $t_{0.10} = 1.310$, $t_{0.05} = 1.697$, $t_{0.025} = 2.042$, $t_{0.01} = 2.457$ 和 $t_{0.005} = 2.750$ 。

我们现在利用 t 分布给出总体均值置信区间的形式。

- 总体均值的置信区间 (confidence intervals for the population mean) (总体方差未知 (population variance unknown)) —— t 分布 (t-distribution)。如果我们从一个方差未知的总体中抽样并且满足下面两个条件之一时,
 - 样本是大样本;
 - 或者样本是小样本但是总体服从正态分布, 或者近似正态分布。

那么, 总体均值 μ 的一个 $(1 - \alpha)\%$ 置信区间为:

$$\bar{X} \pm t_{\alpha/2} \frac{s}{\sqrt{n}} \tag{6-6}$$

式中, $t_{\alpha/2}$ 的自由度数数值为 $n - 1$, n 是样本容量。

例 6-5 再次利用了例 6-4 中的数据, 但是, 该例将利用 t 统计量而不是 z 统计量来计算夏普比率总体均值的置信区间。

例 6-5 夏普比率总体均值的置信区间—— t 分布

如同例 6-4 中所述, 一个投资分析师试图计算一个基于 100 个美国股权共同基金的随机样本的夏普比率总体均值的 90% 置信区间。夏普比率的样本均值为 0.45, 而夏普比率的样本标准差为 0.30。现在夏普比率分布的总体方差是未知的, 该分析师决定利用理论上正确的 t 统计量来计算置信区间。

因为样本容量是 100, 故 $df = 99$ 。在书后的表格中, 最接近的值是 $df = 100$ 。利用 $df = 100$ 并向下查到 0.05 一栏, 我们找到 $t_{0.05} = 1.66$ 。这个可靠性因子比例 6-4 中所用的 $z_{0.05} = 1.65$ 要稍大一些。于是置信区间等于

$$\bar{X} \pm t_{0.05} \frac{s}{\sqrt{n}} = 0.45 \pm 1.66 \frac{0.30}{\sqrt{100}} = 0.45 \pm 1.66(0.03) = 0.45 \pm 0.0498$$

置信区间为 0.4002 ~ 0.4998, 或者保留到小数点后两位的 0.40 ~ 0.50。对于保留到小数点后两位的结果来说, 这个置信区间和例 6-4 中所计算的结果没有差别。

表 6-3 汇总了我们已经使用的不同的可靠性因子。

表 6-3 计算可靠性因子的基础

样本来自于:	小样本的统计量	大样本的统计量	样本来自于:	小样本的统计量	大样本的统计量
方差已知的正态分布	z	z	方差已知的非正态分布	未获得	z
方差未知的正态分布	t	$t^{①}$	方差未知的非正态分布	未获得	$t^{①}$

①使用 z 也是可以接受的。

⊖ $t_{0.10}$, $t_{0.05}$, $t_{0.025}$, $t_{0.01}$ 和 $t_{0.005}$ 这几个数值同样也用来表示 t 分布在给定具体的自由度数时, 在 0.10, 0.05, 0.025, 0.01 和 0.005 显著性水平下的单边临界值。

6.4.3 样本量的选择

什么样的选择会影响一个置信区间的宽度？到目前为止我们的讨论有两个影响宽度的因素：统计量的选择（ t 或者 z ）以及置信度的选择（影响我们选择的 t 或者 z 的特定值）。这两个选择决定了可靠性因子。（回忆一下，一个置信区间具有点估计 $\pm$ 可靠性因子 $\times$ 标准误的结构。）

样本量的选择也会影响一个置信区间的宽度。在所有其他条件不变的情况下，一个更大的样本容量将会缩小置信区间的宽度。回忆一下样本均值标准误的表述：

$$\text{样本均值的标准误} = \frac{\text{样本标准差}}{\sqrt{\text{样本容量}}}$$

我们看到标准误与样本容量平方根是成反比的。随着我们增加样本量，标准误会减少，于是置信区间的宽度同样会缩小。样本量越大，我们能估计的总体参数的精度也越大。^①在所有其他条件相同的情况下，样本越大越好（在那种意义下）。然而在实际应用中，对于增加样本量我们有两点需要去考虑。首先，正如我们在关于夏普比率的例 6-2 中所看到的，样本容量的增加可能导致从多个总体中进行抽样。其次，样本量的增加可能涉及额外的花费，而这可能超过由此获得的额外精度给我们带来的价值。因此，分析师应该在选择样本容量时权衡这 3 个问题：所需的精度、从多个总体中抽样的风险以及取得不同样本量的花费。

例 6-6 短期资本管理者估计客户的净现金流入

一个短期资本管理者希望获得未来 6 个月他所管理客户的资金流入和流出的 95% 的置信区间。他对一个 10 个客户的随机样本进行电话访问，询问他们对于基金增加投入和取出的计划。然后管理人计算以放在管理人处总基金百分比变化作为样本的每个客户现金流的变化。一个正的百分比变化值是指客户账户的净现金流入，而一个负的百分比变化是指客户账户的净现金流出。管理人通过样本中每个账户的相对大小来对每个客户的回答进行权衡，然后计算一个加权平均数。

作为这个过程的结果，短期资本管理人计算出的加权平均值为 5.5%。因此，一个点估计的结果为所有在管理之下的基金数目将在接下去 6 个月增长 5.5%。样本中观测值的标准差是 10%。历史数据的直方图看上去与正态分布非常接近，因此，管理者假设总体是正态的。

(1) 计算总体均值的 95% 的置信区间，并解释你的发现。

管理人决定了解一下如果他使用样本量为 20 或者 30 的相同均值 (5.5%) 和标准差 (10%) 的样本会使得置信区间变成什么样的情况。

(2) 利用 5.5% 的样本均值和 10% 的标准差，计算样本量为 20 和 30 的置信区间。对于样本量为 30 的置信区间利用式 (6-6)。

(3) 解释你在第 1 问和第 2 问所获得的结果。

解：

(1) 因为总体是未知的，而且样本量较小，管理人必须使用式 (6-6) 中的 t 统计量来计算

① 确定所需的样本量以获得希望置信区间的宽度的计算公式是存在的。定义 $E = \text{可靠性因子} \times \text{标准误}$ 。 E 越小，置信区间的宽度越小，因为 $2E$ 是置信区间的宽度。在给定置信度 $(1 - \alpha)$ 获得希望的 E 值的样本量的计算公式为 $n = [(t_{\alpha/2}s)/E]^2$ 。

置信区间。根据样本量 10，得到 $df = 10 - 1 = 9$ 。对于一个 95% 的置信区间来说，他需要使用 $t_{0.025}$ 这个值来对应 $df = 9$ 。根据书后的附表，这个值为 2.262。因此，总体均值的 95% 的置信区间为

$$\bar{X} \pm t_{0.025} \frac{s}{\sqrt{n}} = 5.5\% \pm 2.262 \frac{10\%}{\sqrt{10}} = 5.5\% \pm 2.262(3.162) = 5.5\% \pm 7.15\%$$

总体均值的置信区间从 -1.65% 到 +12.65%。[⊖] 管理人有 95% 的自信认为这个区间包含了总体均值。

(2) 表 6-4 给出了对于这三个样本容量的计算结果。

表 6-4 三个样本容量的 95% 置信区间

分布	95% 置信区间	下界 (%)	上界 (%)	相对大小 (%)
$t(n = 10)$	$5.5\% \pm 2.262(3.162)$	-1.65	12.65	100.0
$t(n = 20)$	$5.5\% \pm 2.093(2.236)$	0.82	10.18	65.5
$t(n = 30)$	$5.5\% \pm 2.045(1.826)$	1.77	9.23	52.2

(3) 置信区间的宽度随着我们样本容量的增加而缩小。这一缩小是关于标准误的一个函数，它会随着 n 的增加而减少。可靠性因子同样也随着自由度数值的增加而变小了。表 6-4 最后一列显示的是对于 $n = 10$ 置信区间宽度的相对大小，在 $n = 10$ 时，该值为 100%。利用样本量为 20 的样本将会缩减置信区间的宽度到 $n = 10$ 置信区间宽度的 65.5%。利用样本量为 30 的样本将会缩减几乎一半的置信区间宽度。比较这些选择，短期资本管理人将利用样本量为 30 的样本获得最精确的结果。

在介绍了抽样和参数估计的大部分基本概念之后，我们将能够把关注焦点放在分析师特别关注的一些抽样问题上。统计推断的质量取决于数据的质量，也取决于所使用抽样计划的好坏。金融数据会产生特殊的问题。另外，抽样计划经常会存在一个或多个偏见。本章的下一节将讨论这些问题。

6.5 抽样中的若干问题

我们已经看到样本区间长度的选择可能产生从多个总体中抽样的问题。^{*}事实上，在处理金融数据过程中，如何做到有效抽样，是存在许多问题的。在本节中我们将讨论 4 个与抽样相关的问题：数据挖掘的偏差，样本选择的偏差，前视偏差和时期偏差。所有这些问题在点估计、区间估计和假设检验中都是非常重要的。正如我们将会看到的，如果样本以任何方式存在偏差，那么由我们所抽取的样本所得到的点估计和区间估计或者是其他任何结论都将是错误的。

6.5.1 数据挖掘的偏差

数据挖掘涉及对于相同或者相关数据的过度使用，我们将马上对于这种过度使用方式进行介绍。数据挖掘偏差是指由于错误使用数据而产生的误差。具有数据挖掘偏差的投资策略通常是不

[⊖] 我们假设在这个例子中的样本容量与客户总体的大小相比是充分得小，所以我们丢弃了有限总体修正因子。

会在将来获得成功的。即使如此，还是有不少投资从业者和研究者经常会从事数据挖掘的工作。因此，分析师需要理解和防范这类问题。

数据挖掘 (data-mining) 是通过广泛地在数据集中搜寻统计上显著的模式 (即重复地“测试”相同的数据直到找到一些看上去起作用的结果为止) 来确定模型的一种行为。^①在涉及统计显著性的操作中，我们设定一个显著性水平，该水平是当假设事实上是正确的，而我们拒绝所检验假设的概率。^②因为我们不希望拒绝一个实际为真的假设，所以检验者经常将显著性水平设定为一个相对较小的数值，如 0.05 或者 5%。^③假设我们检验的假设为一个变量没有准确预测股票收益率，并且我们最终检验了 100 个不同变量。另外，让我们假设在实际上这 100 个变量中没有一个具有预测股票收益率的能力。利用我们检验中的 5% 显著性水平，由于有随机性的影响，我们仍将预计 100 个变量中有 5 个会表现出对于股票收益率的显著预测能力。我们已经挖掘了数据以发现一些表现显著的变量。本质上，在我们得到一些数据集的事后结果时，我们已经对于相同的数据进行了不断的检验。在这个意义上，它就涉及了过度使用数据的数据挖掘。如果我们只报告了显著的变量，而没有报告我们所测试的变量总数，这样是不能很好地作为一个预测量的，我们将提供一个非常误导的结论。我们的结果将看上去比它们实际的情况更加显著，因为根据统计推断理论，如刚才描述的一系列的检验，使得传统对于给定显著性水平 (例如 5%) 的解释无效。

我们应该如何来检验存在的数据挖掘偏差呢？对于大部分的金融数据，最快的方法就是对于所提出的变量或策略做样本外检验。一个样本外检验 (out-of-sample test) 利用与原来变量、策略或模型所使用的样本不相交的时间区间的样本进行检验。如果一个变量或投资策略是数据挖掘的结果，它一般应该不会在样本外的检验中显著。一个在样本外检验中统计上和经济上都显著且具有说服力的经济基础的变量或者投资策略可能是一个有效投资策略的建立基础。然而，仍要对此提高警惕。最重要的样本外检验是未来投资的成功。如果这个策略为其他投资者所知，价格就可能发生调整，从而使得经过良好检验的策略在未来失效。总之，分析师应该注意许多明显有利可图的投资策略可能带有数据挖掘偏差的，因此应该对于公布的投资研究结果在未来谨慎地予以应用。

了解数据挖掘的程度可能是复杂的。为了评估一个投资策略的显著性，我们需要知道有多少不成功的策略被当前的检验者或者以前的检验者用相同或相关的数据集所尝试过。在实际中，许多研究密切地建立在其他检验者所做的基础之上，因此，用麦奎因 (McQueen) 和索利 (Thorley) (1999) 的术语来说，反映了代际间的数据挖掘。代际间的数据挖掘 (intergenerational data mining) 涉及利用以往研究者的结果，并利用一个相同或相关的数据集来指导当前的研究。^④分析师们已经积累了许多对于金融数据集特性的观测结果，并且其他一些分析师们可能在基于他们对于先前其他研究者的研究经验所得到的数据集中，建立很可能会得到支持的模型或者投资策略。结果，那些新结果的重要性可能被夸大了。研究已经表明，这类数据挖掘偏差的程度可能相当

① 一些研究者使用术语“数据窥探”来替代表述数据挖掘。

② 为了加深对于数据挖掘的理解，引进一些与假设相关的检验对我们是非常有帮助的。在假设检验一章中将包含对于显著性水平和显著性检验的深入讨论。

③ 就我们先前讨论的置信区间来说，5% 的显著性水平对应于基于一个合适的样本统计量 (如当假设是关于总体均值时的样本均值) 的总体统计量的假定值落在 95% 置信区间之外。

④ 术语“代际间”来源于将每项研究者的进展视为一代。Campbell、Lo 和 MacKinlay (1997) 已经将代际间数据挖掘称为“数据窥探”。然而后面那个词汇一般常用来作为数据挖掘的同义词；因此，McQueen 和 Thorley 的术语是更加确切的。当我们想要强调所参考的是一个检验者新的或者是独立的数据挖掘时，我们就有术语“代际内数据挖掘”。

之高。^①

在上述定义和解释的背景之下，我们就能够理解麦奎因和索利（1990）基于流行的 Motley Fool “Foolish Four” 投资策略的使人信服的对于数据挖掘的研究。Foolish Four 策略最早于1996年被提出，它是道氏红利策略的一个版本，它被后人修改，能够得到一个比道氏红利策略在1973~1993年更加高的算术平均收益率。^②从1973~1993年，Foolish Four 投资组合具有平均25%的年收益率，并且由报刊声称这个策略在未来也会产生类似的收益率。然而，如同麦奎因和索利所讨论的，Foolish Four 策略受到非常严重数据挖掘偏差的影响，包括从代际间数据挖掘的影响，因为策略的制定者们利用了先前研究者们所使用的数据集中的观测结果。麦奎因和索利通过进一步改进 Foolish Four 投资组合来强调了数据挖掘问题的严重性。他们通过挖掘数据构造了一个 “Fractured Four” 投资组合，它在1973~1996年间获得了几乎35%的收益率，打败了 Foolish Four 策略几乎8个百分点。注意到所有的 Foolish Four 股票在偶数年都表现很好，但是在奇数年表现较差，并且第二低价的高收益股票是相对在奇数年表现最好的股票，于是 Fractured Four 投资组合在偶数年以等权重持有 Foolish Four 股票，而在奇数年只持有第二低价的股票。偶数年 and 奇数年表现的差别有多少可能性反映了对应的经济影响，而不是一个碰巧在这个时间段上出现的数据模式？非常可能的情况是这一点也不可能。除非一个投资策略反映了对应经济的影响，我们一般不会预计投资策略具有预测未来值的能力。因为 Foolish Four 策略也掺杂了数据挖掘问题，同样的问题也会发生在其身上。麦奎因和索利发现在1949~1972年间的样本外检验中，Foolish Four 策略具有与购买并持有 DJIA 大约相同的平均收益率，但是其具有更高的风险。如果考虑 Foolish Four 策略的更高税收和交易成本，这个比较会更加对其不利。

麦奎因和索利指出两个可以提醒分析师们可能存在潜在数据挖掘的迹象：

过多的挖掘/过小的自信度（too much digging/too little confidence）。研究者对于许多变量的检验是数据挖掘问题的“过多挖掘”的警告信号。不幸的是，许多研究者不披露在建立模型中所检验变量的个数。虽然所检验变量的个数可能没有报告，但是我们应该仔细从其字里行间了解到该研究者尝试了许多变量。对于诸如“我们注意到”或者“某人注意到”这类关于一个数据集运行模式的术语表述，应该引起我们对研究者基于他们自身的或者其他人的数据观测结果，而对变量进行多次尝试的怀疑。

没有故事/没有未来（no story/no future）。对于一个变量或者交易策略缺乏清晰的经济学原理解释是数据挖掘问题的“没有故事”的警告信号。对于变量为什么会有效缺乏一个具有说服力的经济学解释或者故事的模型是不可能具有预测能力的。在一个国际金融数据库中利用广泛的变量搜索所得到的结果中，雷恩韦伯（Leinweber）（1997）发现某个远离美国的国家的黄油产量解释了75%的标准普尔500指数中所代表的美国股票收益率的变动情况。这种模式没有任何具有说服力的经济学解释，非常可能是由一个特定时间段的随机模式所产生的。^③如果我们确实对于一个显著的变量有一个具有说服力的经济学解释，那又会怎么样呢？麦奎因和索利提醒我们，一个具有说服力的经济学解释是必要条件，但不是一个交易策略具有价值的充分条件。正如我们之前所提

① 例如，Lo 和 MacKinlay（1990）得到这类关于资本资产定价模型检验的偏差程度是相当大的结论。

② 道氏红利策略，即熟知的道氏狗股策略（Dogs of Dow Strategy），其包含在年初持有由等权重的10个最高收益率的 DJIA 股票组成的投资组合。在 McQueen 和 Thorley 研究的时候，Foolish Four 策略表述如下：在每年年初，Foolish Four 投资组合为从10个最高收益率 DJIA 股票中的5个最低价的股票中购买一个由4只股票组成的投资组合。第5个低价股票被排除在外，然后40%投资于第二低价的股票，剩下的3个各投资20%。

③ 在金融学文献中，这类随机但与未来无关的模式有时被称为数据集的伪造。

到的,如果策略被公之于众,市场价格可能调整以反映新的信息,因为交易者们都试图利用该策略;结果该策略可能会不再起作用。

6.5.2 样本选择的偏差

当研究者们考察分析师或者投资组合经理所关心的问题时,他们可能会根据不同理由(可能因为数据的可获得性),从分析中排除特定的股票、债券、投资组合或者时间段。当数据的可获得性导致特定的资产被排除到分析之外时,我们将由此产生的问题称为样本选择偏差(sample selection bias)。例如,你可能从一个跟踪当前存在于市场中的公司的数据库中进行抽样。比如许多共同基金数据库就只提供那些当前存在基金的历史信息。报告历史资产负债表和利润表信息的数据库遭受到与共同基金数据库同样类型的偏差问题:不再经营的基金或者公司将不会在数据库中出现。因此,利用这类数据库的研究将遭受到样本选择偏差(也被称为生存偏差(survivorship bias))的影响。

蒂姆森(Dimson)、马什(Marsh)和斯汤顿(Staunton)(2002)在国际指数中提出生存偏差的问题:

一个突出的问题是市场生存者对于长期收益率估计的影响。市场不仅仅经历令人失望的业绩表现,同样也通过没收、恶性通货膨胀、国有化和市场失灵而遭受价值上的总损失。通过度量经过长时间段生存下来的市场业绩表现,我们所得到的推断是基于生存者的。但是,正如布朗(Brown)、哥兹曼(Goetzmann)和罗斯(Ross)(1995)以及乔瑞(Jorion)和哥兹曼(1999)所指出的,一个人不能够事先确定哪个市场会生存,哪个市场会毁灭。

当我们同时利用股票价格和会计数据时,生存偏差有时会表现出来。例如,许多金融学的研究使用公司市场价格与每股账面权益的比率(即市净率 P/B),并发现 P/B 是与公司收益率成反比关系的(参见法马(Fama)和弗兰茨(French)(1992, 1993))。 P/B 也被用来构造许多流行的价值型和成长型指数。然而,如果我们所用来收集会计数据的数据库排除了失败的公司,那么就会导致一个生存偏差的结果。科萨里(Kothari)、山肯(Shanken)和斯隆(Sloan)(1995)恰恰检验了这个问题,并认为失败公司的股票预计将具有低收益率和低 P/B 。如果我们排除这些失败公司的股票,那么那些包含在内的具有低 P/B 的股票比将所有低 P/B 股票包含在内的指数,平均地具有更高的收益率。科萨里、山肯和斯隆认为这一偏差是造成先前平均收益率和 P/B 呈反比关系结果的主要原因。^①我们这时可以提供的唯一建议是意识到任何在样本中可能存在的固有偏差。显然,样本选择偏差会给任一研究的结果埋下阴影。

一个样本也可能因为一个公司的股票从交易所中去除(或者摘牌、退市)而产生偏差。^②例如芝加哥大学证券价格研究中心是一个主要的学术研究用收益率数据的提供者。当一个摘牌发生时,CRSP试图去收集退市公司的收益率数据,但是在很多情况下,它无法这么去做,因为其中存在许多困难;CRSP必须简单地将退市公司收益率标为缺损值。沙姆维(Shumway)和维特(Warther)(1999)在《金融学》杂志(*Journal of Finance*)上的一项研究记录了由退市所造成的CRSP纳斯达克收益率数据的偏差。作者表明与较差公司业绩表现相关的退市(如破产)比那些

① 关于检验中数据窥探和生存偏差的讨论,参见Fama和French(1996, p. 80)。

② 摘牌、退市是由于许多原因造成的:兼并,破产,清算或者转移到另一个交易所。

较好或中等公司业绩表现的退市（如兼并或转移到其他交易所）更加普遍。另外，对于小公司来说，退市的发生更加频繁。

即使在数据相当连贯且质量很好的市场，样本选择偏差仍旧会发生。较新的资产类，如对冲基金，甚至可能产生更严重的样本选择偏差。对冲基金是由一些利用不同类别的投资工具进行投资的机构，这种投资模式使得它们免于（难以）监管。一般地，对冲基金不要求公开披露其业绩表现（相反，例如共同基金就需要进行公开披露）。对冲基金自己决定是否想要包含在多个不同对冲基金业绩表现数据库中的一个。具有较差历史记录的对冲基金显然不希望使它们的记录公开，从而产生对冲基金数据库的自选择偏差问题。进一步地，如同冯（Fung）和谢（Hsieh）（2002）所指出的，因为只有具有良好记录的对冲基金才会自愿进入一个数据库中，所以在通常情况下，总体的对冲基金历史业绩表现倾向于比实际情况要好。而且，许多对冲基金数据库都会丢弃那些破产的基金数据，从而在这些数据库中产生生存偏差。甚至那些没有丢弃破产基金数据并试图消除生存偏差的数据集，也仍会存在这个问题，因为有些对冲基金会因为业绩差而不公布其业绩表现。^①

6.5.3 前视偏差

如果一个检验设计使用到了在检验日无法获得的信息，那么该检验设计就受到了前视偏差（look-ahead bias）的影响。例如，使用股票市场收益率和资产负债表会计数据来进行交易策略的检验就必须对于前视偏差进行说明。在这类检验中，一家公司每股账面价值通常被用来构建 P/B（市净率）变量。虽然股票的市场价格是在相同时点为每个市场参与者所及时获取的，但是期末财务报告中的每股账面权益可能在下一季度公布之前不会为公众所获知。

6.5.4 时期偏差

如果一个检验设计是基于一个使得结论成立的特定时间段的话，那么该检验设计就受到了时期偏差（time-period bias）的影响。一个较短的时间序列很可能得出一个无法反映长期的特殊区间段的结果。一个较长的时间序列可能给出实际投资业绩表现的一个更加准确的结果；但是长时间序列的缺点在于在其时间区间内可能产生一个结构性的变化，从而造成其数据来自两个不同的收益率分布。在这种情况下，在变化之前的分布所反映的情况与变化之后的分布所描述的情况会有不同。

例 6-7 投资研究中的偏差

一位分析师正在复核由美国股票历史收益率所得到的实证结果。她发现价值型的股票（即那些低 P/B（市净率）的股票）在最近的某些时间段的表现要优于成长型的股票（即那些高 P/B（市净率）的股票）。在复核了美国市场的结果之后，这位分析师想知道价值型的股票在英国是否也具有这种吸引力。她检验了 1990 年 1 月至 2003 年 12 月英国市场价值型股票和成长型股票的业绩表现。在进行这项研究中，这位分析师做了如下工作：

- 获得当前《金融时报》股票交易所（FTSE）全部股票指数（这是一个市值加权的指数）中的成分股。
- 删除那些没有 12 月年底财务报告的公司。

① 更多关于对冲基金业绩解释的问题，参见 Ackerman, McEnally 和 Ravenscraft（1999）和 Fung 和 Hsieh（2002）。注意，如果一个成功的基金因为它们不再需要新的现金流入而不报告它们的业绩情况时，一个相反类别的偏差也可能出现。

- 使用年末账面价值和市场价格来对余下的那些公司按年末 P/B 进行排序。
- 基于这些排序, 将所有公司分成 10 个投资组合, 每个投资组合包含相同数目的股票。
- 计算在每个排序的时点之后 12 个月的每个投资组合的等权重收益率以及 FTSE 所有股票指数的收益率。
- 将每个投资组合收益率减去 FTSE 收益率, 得到每个投资组合的超额收益率。

根据该分析师的研究设计, 描述并讨论下列每个偏差:

- 生存偏差
- 前视偏差
- 时期偏差

生存偏差。如果一个检验设计无法解释那些破产、兼并或者由于其他原因而被数据集所排除在外的公司, 那么该检验设计就受到了生存偏差的影响。在这个例子中, 该分析师使用了当前 FTSE 中所列出的股票而不是在每年年初时实际在列表中所存在的股票。在一定程度上, 收益率的计算排除了那些指数所移除的公司, 所以具有低 P/B (市净率) 的投资组合的业绩表现遭受了生存偏差, 它可能被高估了。在检验期的一些时间, 当前已经不存在的那些公司被检验所排除了。它们很可能具有较低的价格 (和较低的 P/B) 和较差的收益率。

前视偏差。如果一个检验设计使用到了在检验日无法获得的信息, 那么该检验设计就受到了前视偏差的影响。在这个例子中, 分析师所进行的检验假设了必要的会计信息都能够每个财政年度末获得。例如, 该分析师假设 1990 财政年度的每股账面价值可以在 1990 年 12 月 31 日获得。因为这个信息要到财政年度结束之后的几个月才会披露, 所以这个检验可能存在前视偏差。这个偏差将使得基于该信息的策略看上去是成功的, 但是它实际上假设了完全的预测能力。

时期偏差。如果一个检验设计是基于一个使得结论成立的特定时间段的话, 那么该检验设计就受到了时期偏差的影响。虽然该检验覆盖了一个超过 10 年的时间区间, 但是那个区间可能对于检验一个市场异象来说还是太短了。理想的是, 一位分析师应该在多个商业周期内对于市场异象进行检验, 以保证结论不受到时期偏差的影响。如果所选的时间段对于该策略有利, 那么这一偏差可能赞同其所提出的这项策略。

假设检验

7.1 引言

分析师经常会遇到一些有关金融市场如何运行的相互对立的观点。在这些观点中，一些观点来自人们对于市场的个人研究或者经历，另一些观点来源于与同事间的交流，而这些观点更多的是出现在金融和投资的相关专业文献之中。那么，通常一位分析师是如何来判断一个关于金融市场的陈述可能是正确的还是可能是错误的呢？

当我们能够将一个想法或者论断浓缩为关于一个量（如对应的均值或者总体均值）的确定性表述时，这个想法就成为了一个在统计上可以检验的陈述或者假设。分析师可能会希望回答下列问题：

- 这个基金所对应的平均收益率是否不同于其基准的平均收益率？
- 这只股票收益率的波动率是否在其加入指数之后发生了改变？
- 一个证券的买卖差价是否与市场中对该证券进行做市的交易商的数量有关？
- 从一个国家的债券市场中所获得的数据是否支持经济理论对于利率期限结构（收益率和到期期限之间的关系）的预测？

为了回答这些问题，我们需要使用假设检验的概念和工具。假设检验是统计推断（基于实际观测到的小部分数据（一个样本）对于庞大数据（一个总体）作出判断的过程）中的一个主题。假设检验的概念和工具提供了判断所获得的证据是否能够支持假设的一种客观方法。虽然我们的结论总是以确定性的语言进行表述的，但是在对于一个假设进行统计检验之后，我们应该能对该假设是否为真的概率有一个比较清晰的认识。假设检验已经成为投资这一门学科发展中的一个有用的工具。正如社会研究协会的罗伯特 L. 卡恩（Robert L. Kahn）所说的，“只有当假设和数据相互紧密地联系在一起时，科学这个磨坊中才会研磨出有用的产品。”

本章的主要关注点在于假设检验的框架和关于均值和方差（这两个经常在投资中使用的量）的检验。我们在下一节中将给出一个假设检验过程的概貌。然后，我们将介绍关于单个均值的假

设检验和关于均值间差异的假设检验。在本章的第4节中，我们将介绍关于单个方差的假设检验和对于方差间差别的假设检验。我们在本章的最后将简要介绍一些统计推断中的其他重要问题和使用的技术。

7.2 假设检验

正如我们所提到的，假设检验是统计学分支——统计推断的一个组成部分。传统上，统计推断的研究领域具有两个子问题：参数估计（estimation）和假设检验（hypothesis testing）。参数估计所回答的问题是“参数值（例如总体均值）到底是多少？”该问题是用包含点估计的一个置信区间来回答的。拿总体均值的例子来说：我们使用包含以点估计获得的样本均值来对于总体均值建立一个置信区间。为了具体说明这个问题，假设样本均值是50，一个总体均值的95%置信区间为 50 ± 10 （置信区间从40~60）。如果我们正确地构造出了这个置信区间，那么40~60的区间包含总体均值的概率为95%。^①第2个统计推断的分支，假设检验具有稍稍不同的关注点。一个假设检验回答的问题是“参数的值（比如说总体均值）是否是45（或者是其他的具体值）？”“总体均值是45”这个断言是一个假设。一个假设（hypothesis）被定义为关于一个或者多个总体的论断。

本节所关注的是假设检验的概念，假设检验的过程是获得信息的严格方法（或科学方法）中的一部分。该科学方法从观测开始，逐步建立一套组织和解释观测值的理论。我们通过一个理论做出正确预测的能力来判断该理论的正确与否。例如，对于新的观测值出现结果的预测。^②如果预测是正确的，我们将继续把该理论作为一个对于我们观测值的可能的正确解释。当风险在观测的结果中起作用时，正如在金融中我们所面对的，我们可以只试图做出对于新的数据是否支持我们预测的无偏的基于概率的判断。统计假设检验在对不确定环境中假设的检验中扮演了重要的角色。在一位分析师每天的工作中，他可能会以不同质量的回答来应对各种问题。当他将问题正确地转化为一个可检验的假设并通过假设检验做出报告时，他已经提供了一个与科学方法的标准相一致的支持其假设的要素。当然，分析师的逻辑推理、经济解释、信息来源以及其他可能的因素同样在我们对于回答质量的评估中起到了重要作用。^③

我们将假设检验围绕着下面7个步骤进行介绍。

- **假设检验的步骤。**检验一个假设的步骤如下：^④
 1. 给出假设。
 2. 找到适当的检验统计量及其概率分布。
 3. 确定显著性水平。
 4. 给出决策规则。
 5. 收集数据并计算检验统计量。
 6. 做出统计推断。
 7. 做出经济决策或者是投资决策。

① 我们在抽样一章中对于置信区间的建立和解释已经进行了讨论。

② 为了可以被检验，一个理论必须具有预测错误结果的能力。

③ 参见 Freely 和 Steinberg（1999）对于理性决策批判性思考的讨论。

④ 这个步骤列表是基于 Daniel 和 Terrell（1986）所得到的。

我们将利用一个关于加拿大股票风险溢价符号的假设检验的例子来介绍并解释每个步骤。上面的步骤建立了一个假设检验的传统方法。我们将在本节结尾介绍一个常用的对于这些步骤的替代方法——p 值方法。

假设检验的第 1 步是给出假设。我们经常给出两个假设：原假设（或零假设，null），记为 H_0 ；备择假设，记为 H_1 。

- 原假设（null hypothesis）的定义。原假设是要被检验的假设。例如，我们可能假设加拿大股票风险溢价的总体均值小于等于 0。

原假设是一个被认为正确的命题，除非我们所用来做假设检验的样本给出令人信服的证据说明该原假设是错误的。当这种证据被给出时，我们就会接受备择假设。

- 备择假设（alternative hypothesis）的定义。备择假设是在原假设被拒绝时所接受的假设。我们这里的备择假设为加拿大股票风险溢价的总体均值大于 0。

假设我们的问题是关于一个总体参数 θ 与一个可能的值 θ_0 之间的关系（这两个参数分别读作“theta”和“theta sub zero”）。[⊖] 总体参数包括总体均值 μ 和总体方差 σ^2 。我们能给出 3 种不同的假设，我们将根据备择假设的论述来给予它们标注不同的类别。

- 假设的给出。我们可以以 3 种不同的方式给出原假设和备择假设：
 1. $H_0: \theta = \theta_0$ 以及 $H_a: \theta \neq \theta_0$ （一个“不等”的备择假设）
 2. $H_0: \theta \leq \theta_0$ 以及 $H_a: \theta > \theta_0$ （一个“大于”的备择假设）
 3. $H_0: \theta \geq \theta_0$ 以及 $H_a: \theta < \theta_0$ （一个“小于”的备择假设）

在我们加拿大股票的例子中， $\theta = \mu_{RP}$ ， μ_{RP} 代表加拿大股票风险溢价的总体均值，而 $\theta_0 = 0$ ，并且这里我们使用的是上面 3 种假设给出方式的第 2 种。

第 1 种假设给出的方式是一个双边的假设检验（two-sided hypothesis test）（或者双侧假设检验（two-tailed hypothesis test））：如果证据显示总体参数或者小于 θ_0 或者大于 θ_0 时，我们就拒绝原假设而接受备择假设。相反，第 2 种和第 3 种假设给出方式都是单边假设检验（one-sided hypothesis test）（或者单侧假设检验（one-tailed hypothesis test））。对于第 2 种和第 3 种情况，只有当证据显示总体参数大于 θ_0 （第 2 种情况）或者小于 θ_0 （第 3 种情况）时，我们才会拒绝原假设。备择假设是单侧的。

注意到在上面每个例子中，我们给出的原假设和备择假设能够表示所有可能的参数值。例如，在第 1 种假设给出的方式中，参数或者等于假设值 θ_0 （在原假设下）或者不等于假设值 θ_0 （在备择假设下）。这两个叙述在逻辑上穷举了参数的可能值。

尽管给出假设的方法不同，但是我们总是可以对取到等式 $\theta = \theta_0$ 这点上的原假设做出检验。无论原假设是 $H_0: \theta = \theta_0$ ， $H_0: \theta \leq \theta_0$ 或者 $H_0: \theta \geq \theta_0$ ，我们在实际上检验的都是 $\theta = \theta_0$ 。这个理由是很直接的。假如参数的假设值为 5。考虑 $\theta \leq 5$ ，具有一个“大于”的备择假设， $H_a: \theta > 5$ 。如果我们具有充分的证据拒绝 $\theta = 5$ 而接受 $H_a: \theta > 5$ ，那么我们就直接具有了拒绝比参数 θ 更小值，如 4.5 或 4 这个假设的充分证据。回顾一下，检验原假设的计算方法对于所有这 3 个假设给出的形式都是相同的。而我们将马上看到，这 3 种情况的不同点在于如何来评价计算的结果以

⊖ 希腊字母，诸如 σ ，都是用来表示总体的参数的；而斜体的罗马字母，例如 s ，都是用来表示样本统计量的。

确定是否拒绝原假设。

我们如何选择原假设和备择假设？我们很可能会使用最常用的“不等于”的备择假设。我们会因为证据显示参数或者大于或者小于 θ_0 ，而拒绝原假设。然而，有时我们可能会具有一个“怀疑”或者“希望”的条件，我们想要找到支持该条件的证据。[⊖]在这种情况下，我们可以将该条件为真的表述作为其备择假设；而我们所检验的原假设表述为该条件不成立。如果证据支持拒绝原假设，而接受备择假设，那么我们在统计意义上证明了我们所认为的为真的命题。例如，经济理论表明投资者要求一个正的对于股票的风险溢价（风险溢价（risk premium）被定义为股票的期望收益率减去无风险利率）。根据将备择假设表述为“所希望的”条件的原则，我们给出下述的假设：

H_0 ：加拿大股票风险溢价的总体均值小于或等于0。

H_a ：加拿大股票风险溢价的总体均值为正。

注意“大于”和“小于”的备择假设反映了研究者的一种比“不等”的备择假设更加强烈的信念。为了强调中立的态度，当一个单边的备择假设同样合理时，研究者有时也许会选择一个“不等”的备择假设。

假设检验的第2步是找到适当的检验统计量及其概率分布。

- 检验统计量（test statistic）的定义。一个检验统计量是一个由样本计算而来的量，其数值是我们决定是否拒绝原假设的一个基础。

我们统计决策的关注点在于检验统计量的值。通常（在本章中我们检验的许多例子中），检验统计量具有如下形式：

$$\text{检验统计量} = \frac{\text{样本统计量} - \text{在 } H_0 \text{ 假设下总体参数的值}}{\text{样本统计量的标准差}} \tag{7-1}$$

对于风险溢价的例子来说，我们所关心的总体参数是风险溢价的总体均值 μ_{RP} 。我们将 H_0 下的总体均值的假设值以 μ_0 表示。使用符号来重新表述该假设为，我们检验 $H_0: \mu_{RP} \leq \mu_0$ 对应于 $H_a: \mu_{RP} > \mu_0$ 。然而，因为在原假设下我们检验 $\mu_0 = 0$ ，所以我们将假设写为： $H_0: \mu_{RP} \leq 0$ 对应于 $H_a: \mu_{RP} > 0$ 。

样本均值提供对于总体均值的一个估计。因此，我们能使用由历史数据 $\bar{X}_{RP}$ 所计算得到的风险溢价的样本均值作为式（7-1）中的样本统计量。样本统计量的标准差被称为统计量的“标准误”，它是式（7-1）中的分母。对于这个例子，样本统计量是一个样本均值。对于一个由标准差为 σ 总体所产生的样本所计算得到的样本均值 $\bar{X}$ ，其标准误由下面两个表达式之一给出：

$$\sigma_{\bar{X}} = \frac{\sigma}{\sqrt{n}} \tag{7-2}$$

当我们已知 σ （总体的标准差）时，或者

$$s_{\bar{X}} = \frac{s}{\sqrt{n}} \tag{7-3}$$

当我们不知道总体标准差，并需要通过样本标准差 s 来估计它时。对于这个例子，因为我们不知道产生收益率过程的总体标准差，所以我们使用式（7-3）。因此，检验统计量为

⊖ 这个对于假设选择讨论中的部分内容是根据 Bowerman 和 O'Connell（1997，p. 386）而来的。

$$\frac{\bar{X}_{RP} - \mu_0}{s_{\bar{X}}} = \frac{\bar{X}_{RP} - 0}{s/\sqrt{n}}$$

在用 0 替代 μ_0 时，我们使用了一个已知的事实：我们检验在等式这点上成立的任何原假设，即这里 $\mu_0 = 0$ 的事实。

我们已经找到了一个用来检验原假设的检验统计量。那么该检验统计量服从什么概率分布呢？我们将在本章中遇到检验统计量的 4 个概率分布：

- t 分布（对于一个 t 检验）
- 标准正态或者 z 分布（对于一个 z 检验）
- 卡方 (χ^2) 分布（对于一个 χ^2 检验）
- F 分布（对于一个 F 分布）

我们将在之后讨论它们的具体细节，但是假设我们能基于中心极限定理来进行一个 z 检验，因为我们的加拿大样本具有许多观测值。[⊖] 总之，对于该风险溢价均值假设检验的检验统计量为 $\frac{\bar{X}_{RP}}{S_{\bar{X}}}$ 。因为我们能有理由地认为该检验统计量服从一个标准正态分布，所以我们可以进行一个 z 检验。

假设检验的第 3 步是确定显著性水平。当检验统计量已经被计算出来时，可能出现两种决策：（1）我们拒绝原假设；（2）我们不能拒绝原假设。我们采取的行为是基于计算出的检验统计量和一个特定可能值之间的比较。我们所选择的比较值是基于所选择的显著性水平。显著性水平反映了我们需要多少样本证据来拒绝原假设。类似于在法庭上所对应的判决，根据假设的性质和所产生错误后果的严重程度，其所要求证明的标准会发生改变。当我们检验一个原假设时，有 4 种可能的结果：

1. 我们拒绝一个错误的原假设。这是一个正确的决策。
2. 我们拒绝一个正确的原假设。这被称为第一类错误（type I error）。
3. 我们不能拒绝一个错误的原假设。这被称为第二类错误（type II error）。
4. 我们不能拒绝一个正确的原假设。这是一个正确的决策。

我们在表 7-1 中列出了这些结果。

当我们在一个假设检验中做出一个决策时，我们冒着犯第一类错误或者犯第二类错误的风险。这两个错误间是互斥的：如果我们错误地拒绝了原假设，我们只可能犯第一类错误；如果我们错误地没有拒绝原假设，我们只可能犯第二类错误。

在检验一个假设中第一类错误的概率记为希腊字母阿尔法 α 。这个概率也被称为检验的显著性水平（level of significance）。例如，检验的一个 0.05 的显著性水平表示有 5% 的概率拒绝一个正确的原假设。第二类错误的概率记为希腊字母贝塔 β 。

控制这两种错误的概率需要我们进行权衡。在所有其他情况相同的条件下，如果我们通过设

表 7-1 假设检验中的第一类错误和第二类错误

决策	真实情况	
	H_0 为真	H_0 为假
不能拒绝 H_0	正确决策	第二类错误
拒绝 H_0 （接受 H_a ）	第一类错误	正确决策

⊖ 中心极限定理认为当样本量很大时，样本均值的抽样分布将近似服从均值为 μ ，方差为 σ^2/n 的正态分布。对于这个例子，我们所使用的样本具有 103 个观测值。

定一个较小的显著性水平（比如 0.01，而不是 0.05）来减少第一类错误的概率，那么我们会增加犯第二类错误的概率，因为我们将不太经常拒绝原假设，包括当它实际为错误的时候。唯一的
同时降低这两种错误概率的方法为增加样本容量 n 。

在实践中给予这两种错误之间的权衡定量通常是不可能的，因为第二类错误本身是难以量化的。考虑 $H_0: \theta \leq 5$ 对应于 $H_a: \theta > 5$ 。因为每个 θ 的真实值大于 5 使得原假设为假，每个 θ 值大于 5 具有一个不同的 β （第二类错误发生的概率）。相反，它对于在 $\theta = 5$ （在该点对于原假设进行检验）的第一类错误发生概率的描述是充分的。因此，当我们进行假设检验时，通常我们只具体地设定 α （第一类错误发生的概率）。但是一个检验的显著性水平是错误拒绝原假设的概率，一个检验的势（power of a test）是指正确拒绝原假设的概率，即在原假设为假时拒绝它的概率。[⊖] 当我们能获得超过一个进行检验的检验统计量时，我们应该在其他条件都相同的情况下，选择势最大的那个检验统计量。[⊗]

总而言之，假设检验的标准方法只涉及显著性水平的确定（第一类错误发生的概率）。在计算检验统计量之前先设定显著性水平最合适。如果我们在检验统计量的计算之后再设定显著性水平，那么我们可能会受到计算结果的影响，从而使我们偏离检验的目标。

我们经常用来进行假设检验的显著性水平有 3 个，分别为：0.10、0.05 和 0.01。定性地来讲，如果我们在 0.10 的显著性水平下拒绝了原假设，那么我们就具有了原假设为假的一些证据。如果我们在 0.05 的显著性水平下拒绝了原假设，那么我们就具有了原假设为假的比较强的证据。而如果我们在 0.01 的显著性水平下拒绝了原假设，那么我们就具有了原假设为假的相当强的证据。对于风险溢价的例子，我们将设定一个 0.05 的显著性水平。

假设检验的第 4 步是表述决策规则。一般的原则可以简单地进行表述。当我们对于原假设进行检验时，如果我们发现计算出的检验统计量的值比在显著性水平 α 下给定的值要偏离很多，那么我们会拒绝原假设。我们认为该结果是统计意义上显著的（statistically significant）。反之，我们将不能拒绝原假设，并认为该结果是统计意义上不显著的（not statistically significant）。我们用来做出决策的，与计算出的检验统计量比较的值是该检验的拒绝点（临界值）。[⊙]

- 检验统计量的拒绝点（rejection point for the test statistic）（临界值（critical value））的定义。一个检验统计量的拒绝点（临界值）是与计算出的检验统计量比较的，以判断是否应该拒绝原假设的值。

对于一个单侧检验，我们利用带有表示确定的第一类错误发生概率 α 的下标的检验统计量来表示一个拒绝点；例如， z_α 。对于双侧检验，我们用 $z_{\alpha/2}$ 表示。为了说明如何使用拒绝点，我们假设使用一个 z 检验，并选择一个 0.05 的显著性水平。

对于一个检验 $H_0: \theta = \theta_0$ 对应于 $H_a: \theta \neq \theta_0$ ，存在两个拒绝点，一个为负，另一个为正。对于一个在 0.05 显著性水平下的双边检验，第一类错误发生的概率必须加总为 0.05。因此，在原假设下，检验统计量分布的拒绝点为 $0.05/2 = 0.025$ 。因此，两个拒绝点为 $z_{0.025} = 1.96$ 以及 $-z_{0.025} = -1.96$ 。令 z 代表检验量所计算出的值。如果我们发现 $z < -1.96$ 或者 $z > 1.96$ ，我们就拒绝原假设。如果 $-1.96 \leq z \leq 1.96$ ，我们就不能拒绝原假设。

对于一个在 0.05 显著性水平下的检验 $H_0: \theta \leq \theta_0$ 对应于 $H_a: \theta > \theta_0$ ，拒绝点为 $z_{0.05} = 1.645$ 。

⊖ 一个检验的势事实上是 1 减去第二类错误发生的概率。

⊗ 然而，我们不总是具有用来比较检验统计量相对检验的势的信息。

⊙ “拒绝点”是对于更加传统的“临界值”术语的一个描述性的同义词。

如果 $z > 1.645$ ，那么我们就拒绝原假设。使得 5% 的结果位于分布右侧尾部的标准正态分布的临界值为 $z_{0.05} = 1.645$ 。

对于一个在 0.05 显著性水平下的检验 $H_0: \theta \geq \theta_0$ 对应于 $H_a: \theta < \theta_0$ ，拒绝点为 $-z_{0.05} = -1.645$ 。如果 $z < -1.645$ ，我们拒绝原假设。

图 7-1 利用 z 检验给出了在 0.05 显著性水平下的 $H_0: \mu = \mu_0$ 对应于 $H_a: \mu \neq \mu_0$ 的检验。“接受域”是对于一组我们无法拒绝原假设检验统计量值的传统名称。（然而，传统名称是不准确的。我们应该避免使用诸如“接受原假设”的词汇，因为这一表述隐含着当我们错误地拒绝原假设时，我们对于原假设的确信程度要高于该检验所保证的程度。[⊖]）在接受域的每一侧都是一个拒绝域（或者临界域）。如果原假设 $\mu = \mu_0$ 是正确的，检验统计量具有 2.5% 的机会落入左侧的拒绝域，或者 2.5% 的机会落入右侧的拒绝域。任何计算出的检验统计量的值落入这两个区间之中的一个，都会造成我们在 0.05 显著性水平下拒绝原假设。拒绝点 1.96 和 -1.96 被视为接受域和拒绝域之间的分界线。

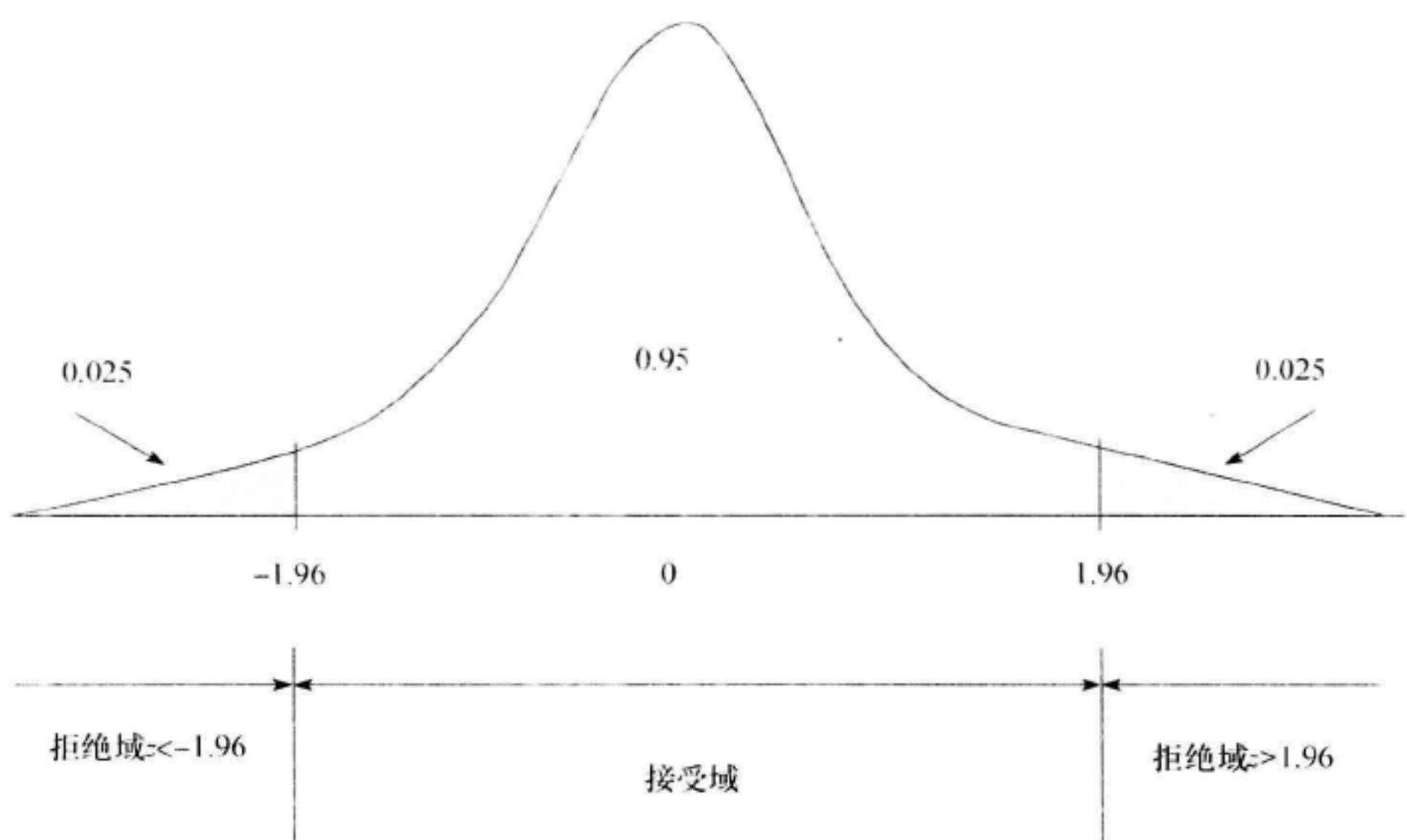


图 7-1 利用 z 检验所得 0.05 显著性水平下的总体均值双边检验的拒绝点（临界值）

图 7-1 给我们提供了一个良好的机会来显示置信区间和假设检验间的关系。一个基于样本均值 $\bar{X}$ 的总体均值 μ 的 95% 的置信区间是由 $\bar{X} - 1.96s_{\bar{X}}$ 到 $\bar{X} + 1.96s_{\bar{X}}$ 区间给出的，其中 $s_{\bar{X}}$ 是样本均值的标准误（式 (7-3)）。[⊖]

现在考虑一个拒绝原假设的条件：

$$\frac{\bar{X} - \mu_0}{s_{\bar{X}}} > 1.96$$

这里， μ_0 是总体均值的假定值。这个条件表示如果检验统计量超过 1.96，就允许拒绝原假设。在等式两边同时乘以 $s_{\bar{X}}$ ，我们有 $\bar{X} - \mu_0 > 1.96s_{\bar{X}}$ ，或者在重新整理后，为 $\bar{X} - 1.96s_{\bar{X}} > \mu_0$ ，而

⊖ 在法庭上（比如说在美国的法庭上）的一个类比为：如果一个陪审团不能给出一个有罪判决（备择假设），那么我们可以认为陪审团不能拒绝原假设，即被告是无罪的，这一说法更加准确。

⊖ 在假设检验中，我们也可以在大样本的情况下，使用这个基于标准正态分布的置信区间。一个备择假设检验和置信区间使用的是 t 分布，而这需要我们在下一节中介绍的概念。

其我们有时也写为 $\mu_0 < \bar{X} - 1.96s_{\bar{X}}$ 。这一表达式说明，如果假定的总体均值 μ_0 小于基于样本均值所得到的 95% 置信区间的下界，那么我们必须 在 5% 显著性水平下拒绝原假设（检验统计量落入拒绝域的右侧）。

现在，我们可以采用其他拒绝原假设的条件：

$$\frac{\bar{X} - \mu_0}{s_{\bar{X}}} < -1.96$$

并利用之前采用的代数方法，将其重写为 $\mu_0 > \bar{X} + 1.96s_{\bar{X}}$ 。如果假定的总体均值要大于 95% 置信区间的上界，那么我们会在 5% 显著性水平下拒绝原假设（检验统计量落入左侧的拒绝域中）。因此，在双边假设检验中一个 α 的显著性水平可以以相同精确的方式表述为一个 $(1 - \alpha)$ 的置信区间。

总之，当原假设条件下的总体参数假定值落在相应置信区间之外时，原假设就会被拒绝。我们可以利用置信区间来检验假设；然而，实务者通常不这么做。计算一个检验统计值（一个数，相比常用置信区间的两个数）会更加地有效率。同样，分析师们只有在很少情况下才会面临单边置信区间的实际问题。而且，只有在我们计算出一个检验统计量时，我们才能获得一个 p 值（一个与我们结果显著性相关的有用的量，我们马上对 p 值进行讨论）。

回到我们风险溢价的检验中，我们给出的假设为 $H_0: \mu_{RP} \leq 0$ 对应于 $H_a: \mu_{RP} > 0$ 。我们找出的检验统计量为 $\frac{\bar{X}_{RP}}{s_{\bar{X}}}$ ，并认为其服从一个标准正态分布。因此，我们在做的是单边 z 检验。我们设定一个 0.05 的显著性水平。对于这个单边 z 检验，0.05 显著性水平下的拒绝点为 1.645。如果计算出的 z 统计量大于 1.645，我们将拒绝原假设。图 7-2 对于该检验进行了描述。

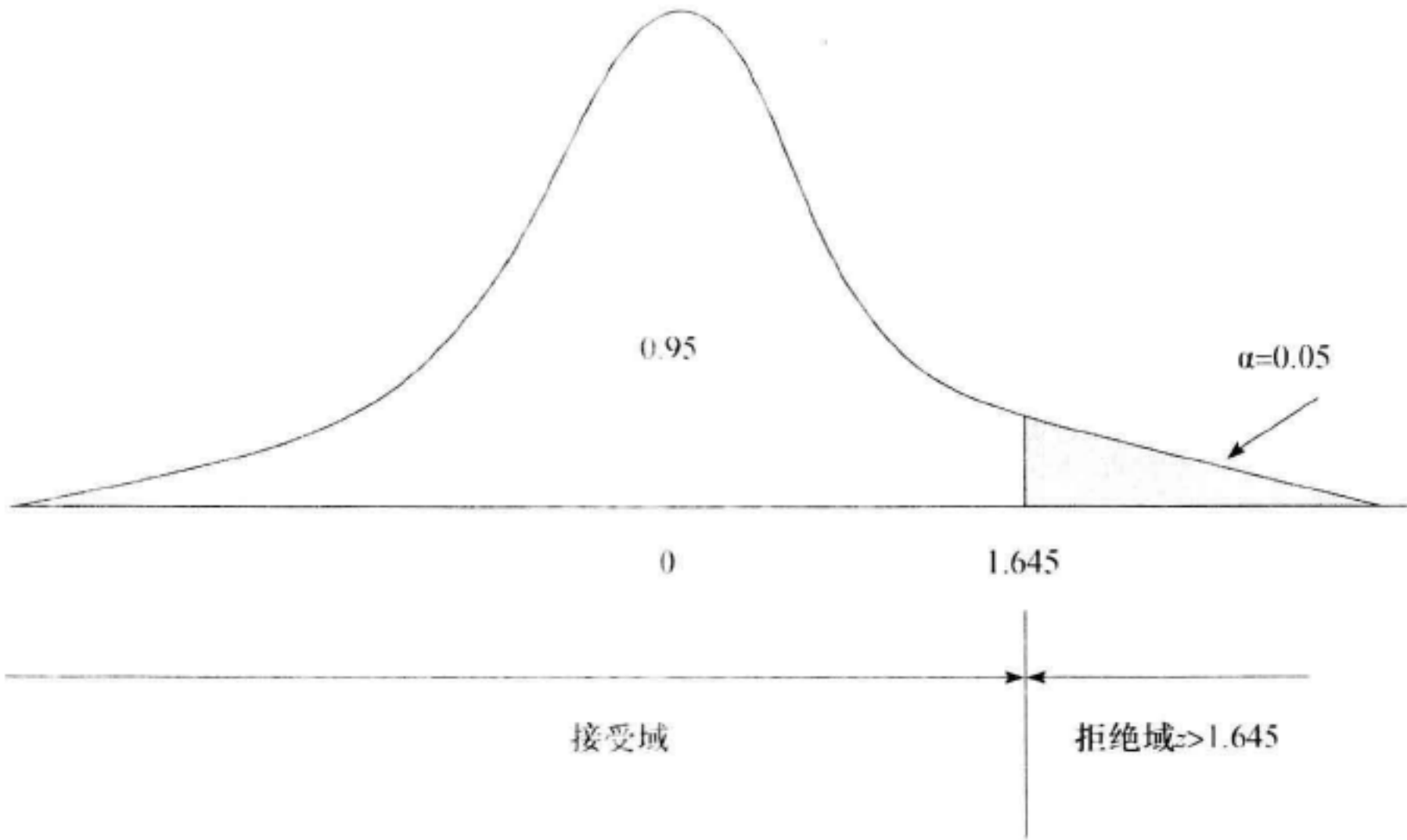


图 7-2 利用 z 检验得到 0.05 显著性水平下的总体均值的单边检验的拒绝点（临界值）

假设检验的第 5 步是收集数据和计算检验统计量。我们结论的质量不仅取决于统计模型的恰当性，还取决于我们在进行检验时所使用的数据的质量。我们首先需要检查所记录的数据的测量误差。其他一些需要注意的问题包括样本选择偏差和时期偏差。样本选择偏差是指由于根据一个特定的属性而系统性地排除总体中的一些元素所造成的偏差。样本选择偏差中的一类为生存偏差。例如，如果我们将样本定义为当前运营的美国债券共同基金，而我们只收集了这些基金的收益率，那么我们会系统性地排除没有经营到今天的那些基金。一般来说，没有生存下来的基金

很可能表现为目前生存下来的表现较差的那些基金；结果样本所反映出的业绩表现可能有一个向上的偏差。时期偏差是指当我们使用一个时间序列样本，我们的统计结论可能对于样本起始日期和终止日期敏感的概率。^①

继续我们关于风险溢价的假设，我们关注加拿大的股票。根据蒂姆森、马什和斯汤顿（2002）更新到2002年年末的数据^②，对于从1900~2002年（103个年度观测值），加拿大股票相对于债券收益率的风险溢价的算术平均值 $\bar{X}_{RP}$ 为每年5.1%。年度风险溢价的样本标准差为18.6%。利用式（7-3），样本均值的标准误为 $s_{\bar{X}} = \frac{s}{\sqrt{n}} = \frac{18.6\%}{\sqrt{103}} = 1.833\%$ 。检验统计量为 $z = \frac{\bar{X}_{RP}}{s_{\bar{X}}} = 5.1\% / 1.833\% = 2.78$ 。

假设检验的第7步，即最后一步是做出经济的或者投资学的解释。经济或者投资学的解释不仅考虑统计的决策，而且要考虑所有相关的经济问题。在第6步中，我们得到了较强的加拿大股票风险溢价为正的证据。所估计的风险溢价的大小为每年5.1%，这在经济上也是非常有意义的。基于这些考虑，一个投资者可能决定参与到投资加拿大股票的基金之中。一些非统计的因素，诸如投资者对于风险的容忍程度以及金融头寸，可能也会在决策过程中加以考虑。

之前的讨论引出了经常在决策过程中出现的问题。我们经常发现一个变量和其假定值之间的微小差别是统计上显著的，但不是经济上显著的。例如，我们可能检验一个投资策略，它基于大样本拒绝了该策略平均收益率为0的原假设。式（7-1）表明在所有其他条件相同的情况下，样本统计量的标准误越小（公式中的除数），检验统计量的值越大，于是原假设会被拒绝的可能性也越大。随着样本大小 n 的增加，标准误会减少，所以对于非常大的样本，我们可能会因为微小偏差而拒绝原假设。我们可能发现虽然一个策略提供一个统计上显著的的正的平均收益率，但是当我们考虑了交易成本、税收和风险之后，结果在经济上就会是不显著的。甚至我们得到一个策略的结果在经济上有意义的结论，我们也应该在实施该策略之前，探究一下为什么该策略会在未来同样起作用的逻辑与理由。这些考虑都不能在假设检验中得到体现。

在结束假设检验过程这个主题之前，我们应该讨论一下一个重要的假设检验中的替代方法，称为 p 值方法。分析师和研究者经常报告假设检验的 p 值（也被称为边际显著性水平）。

- p 值（ p -value）的定义。 p 值是指一个最小的显著性水平，在该水平上，原假设可以被拒绝。

对于风险溢价假设检验中的检验统计量2.78这个值，通过一个标准正态分布的电子表格函数，我们计算出 p 值为0.002 718。我们在此显著性水平下能拒绝原假设。 p 值越小，拒绝原假设、接受备择假设的证据就越强。一个参数等于0的双边假设检验的 p 值经常自动地由统计和计量软件程序产生。^③

我们可以在假设检验的框架下使用上面给出的 p 值作为使用拒绝点的一个替代。如果 p 值比我们设定的显著性水平小，我们就要拒绝原假设。反之，我们不能拒绝原假设。通过这样利用 p 值，我们得到了与之前使用拒绝点相同的结论。例如，因为0.002 718小于0.05，所以我们将拒绝风险溢价检验中的原假设。然而， p 值比拒绝点方法，给我们提供了更加精确的关于证据强度的信息。0.002 718的 p 值表明原假设是在一个远小于0.05显著性的水平上被拒绝的。

如果一个研究者利用0.05显著性水平对于一个问题进行检验，而另一个研究者利用0.01的

① 这些概念在抽样一章中进行了深入的讨论。

② 根据2003年5月19日与作者的通信后更新的数据。

③ 我们也可以使用电子表格来计算 p 值。例如，在微软的Excel中，我们可以使用工作表的函数TTEST，NORMSDIST，CHIDIST和FDIST来分别计算 t 检验、 z 检验、 χ^2 检验和 F 检验的 p 值。

显著性水平,那么读者可能会在比较两个结果时产生困难。这一问题使得人们采用这样的方法,即给出的假设检验的结果中只含有 p 值而忽略了具体设定的显著性水平(步骤3)。统计结果的解释被留给了阅读这项研究的使用者。这一方法有时也被称为假设检验的 p 值方法。^①

7.3 关于均值的假设检验

关于均值的假设检验在实际应用中最为常见。在本节中,我们会讨论一些针对不同类型问题的这类检验。在第一类(在3.1节中讨论)中,我们检验单个总体的均值是否等于(或大于或小于)某个假定值。然后,在3.2节和3.3节中,我们介绍基于两个样本的均值检验问题。观测到的两个样本均值间的差别是否是由于运气或者确实是两个不同的(总体)均值?当我们具有两个相互独立(一个样本的度量数据和另一个样本的度量数据之间不存在关系)的随机样本,那么我们可以应用3.2节的技术。当样本间是相互依赖的,那么应用3.3节的方法是恰当的。^②

7.3.1 对单个均值的检验

一个想要检验所关心的均值或者总体均值的假设的分析师会在绝大部分的情况下采用 t 检验。一个 t 检验(t -test)是利用一个服从 t 分布的统计量(t 统计量)的假设检验。 t 分布是由单个被称为自由度(df)的参数所定义的概率分布。每个自由度定义了这个分布类中的一个分布。 t 分布与标准正态分布的关系比较密切。和标准正态分布一样, t 分布是对称的,且均值为0。然而, t 分布更加伸展:它具有比1大的标准差(与标准正态分布的标准差1相比)^③,其结果远离均值的概率更大(它比标准正态分布具有更厚的尾部)。随着样本容量增大,自由度也会增加,于是分布向两边伸展的程度下降了, t 分布逐渐趋近于标准正态分布。

为什么在本节中将 t 分布作为假设检验关注的一个焦点?在实务中,投资分析师们需要通过计算一个样本方差来估计总体标准差;即总体方差(或标准差)是未知的。对于未知的方差,理论上正确的检验统计量是 t 统计量。那么如果一个正态分布无法描述总体,我们该怎么办呢? t 检验对于中等偏离(除了异常值和严重偏度的情况)正态分布的情况是稳健的(robust)。^④而当我们对于大样本进行检验时,我们会较少对于实际分布偏离正态分布的程度予以关注。在大样本中,根据中心极限定理,无论总体服从什么样的分布,样本均值都是近似服从正态分布的。一般来说,一个包含30个或者更多观测值的样本通常会被作为一个大样本进行处理,而样本量为29个或者更少的样本会被作为小样本进行处理。^⑤

- 对于总体均值假设检验的检验统计量(test statistic for hypothesis tests of the population)

① Davidson 和 MacKinnon (1993) 对于该方法的优点这样说道:“ p 值的方法不需要强迫我们做出关于原假设的决策。如果我们获得一个 p 值,比如0.000 001,我们将几乎肯定想要拒绝原假设。但是如果获得一个 p 值为0.04,或者甚至0.004,那么我们就不会一定会拒绝原假设。我们可能简单地记录下这个结果,将其作为对于原假设怀疑的一个信息,但是它本身并不是结论性的。我们认为这种对于检验统计的有些不可知论的态度(我们以这种态度将它们仅仅作为我们可能会拒绝也可能不会拒绝的一些信息)通常是我们能够采取的最为明智的方法。”

② 当我们想要检验是否来自两个以上总体的总体均值相等的时候,我们使用方差分析(ANOVA)。我们将在相关性和回归分析一章中,介绍ANOVA的常见应用和回归分析。

③ t 分布方差的计算公式为 $df/(df-2)$ 。

④ 参见 Moore 和 McCabe (1998)。如果所需要的概率计算对于假设违背的情况是不敏感的,那么该统计量就被称为稳健的。

⑤ 虽然这一总结是有用的,但是我们要注意,要获得一个服从近似正态抽样分布的样本均值所需要的样本量是依赖于原始总体偏离正态分布的程度的。对于某些总体,“大样本”可能指的是一个样本量远高于30的样本。

mean) (实际情况——总体方差未知 (practical case——population variance unknown))。如果被抽样的总体具有未知的方差,并具有下列两个条件之一:

1. 样本量较大;
2. 样本量较小,但是被抽样的总体服从正态分布,或者近似正态分布。

那么对于单个总体均值 μ 的假设检验的检验统计量为

$$t_{n-1} = \frac{\bar{X} - \mu_0}{s/\sqrt{n}} \quad (7-4)$$

式中, t_{n-1} 表示自由度为 $n-1$ 的 t 统计量 (n 是样本容量); $\bar{X}$ 表示样本均值; μ_0 表示总体均值的假定值; s 表示样本标准差; t 统计量的分母是对于样本均值标准误的一个估计, $s_{\bar{X}} = s/\sqrt{n}$ 。^①

在例 7-1 中,因为样本量较小,该检验被称为对于总体均值的小样本检验。

例 7-1 一个股票型共同基金的风险收益特征 (1)

你现在正在对森达 (Sendar) 股票型基金进行分析,这是一家已经在市场中生存了 24 个月的中等市值成长型基金。在这个区间中,该基金实现了 1.50% 的月度平均收益率,而且该月度收益率的样本标准差为 3.60%。给定该基金所面临的系统性风险 (市场风险) 水平,并根据一个定价模型,我们预期该共同基金在这个区间中应该获得 1.10% 的月度平均收益率。假定收益率服从正态分布,那么实际结果是否和 1.10% 这个理论上的月度平均收益率或者总体月度平均收益率相一致?

- (1) 给出与该研究目的的语言描述相一致的原假设和备择假设。
- (2) 找出对于第 1 问中的假设进行检验的检验统计量。
- (3) 求出 0.10 显著性水平下第 1 问中所检验的假设的拒绝点。
- (4) 确定是否应该在 0.10 显著性水平下拒绝原假设。(利用书后的数值表。)

解:

(1) 我们有一个“不等”的备择假设,其中 μ 是森达股票基金的对应的平均收益率—— $H_0: \mu = 1.10$ 对应于 $H_a: \mu \neq 1.10$ 。

(2) 因为总体方差是未知的,我们利用 $24 - 1 = 23$ 自由度的 t 检验。

(3) 因为这是一个双边检验,我们的拒绝点 $t_{\alpha/2, n-1} = t_{0.05, 23}$ 。在 t 分布的数值表中,我们查看自由度为 23 的行和 0.05 的列,找到 1.714。双边检验的两个拒绝点是 1.714 和 -1.714。如果我们发现 $t > 1.714$ 或者 $t < -1.714$,我们将拒绝原假设。

$$(4) t_{23} = \frac{1.50 - 1.10}{3.60/\sqrt{24}} = \frac{0.40}{0.734847} = 0.544331 \text{ 或者 } 0.544$$

因为 0.544 不满足 $t > 1.714$ 和 $t < -1.714$,所以我们不能拒绝原假设。

置信区间方法提供了关于该假设检验的另一种观点。理论上,基于样本量 n ,服从一个方差未知的正态分布的总体均值的正确的 $100(1 - \alpha)\%$ 的置信区间为

$$\bar{X} - t_{\alpha/2} s_{\bar{X}} \text{ 到 } \bar{X} + t_{\alpha/2} s_{\bar{X}}$$

① 这里需要参考一个技术性的要点。当一个来自于有限总体的样本,无论从式 (7-2) 还是式 (7-3) 所计算得到的均值标准误的估计值都会高于真实的标准误。解释一下,所计算的标准误乘以的是被称为有限总体修正因子的缩减因子,等于 $\sqrt{(N-n)/(N-1)}$,其中, N 是总体容量, n 是样本容量。当样本容量相对于总体容量较小时 (小于 5% 的总体容量), fpc 经常被忽略。与有放回的抽样不同,只有在不放回抽样 (元素在被抽出后不再放回) 的一般情形下才会产生高估的问题。

其中, $t_{\alpha/2}$ 是使得分布右侧尾部的概率为 $\alpha/2$ 的 t 值, 而 $t_{-\alpha/2}$ 是使得分布左侧尾部的概率为 $\alpha/2$ 的 t 值 (对于自由度 $n-1$ 的情况)。这里的 90% 置信区间是从 $1.5 - (1.714)(0.734\ 847) = 0.240$ 到 $1.5 + (1.714)(0.734\ 847) = 2.760$ 的区间, 简写为 $[0.240, 2.760]$ 。平均收益率的假定值 1.10% 落入这个置信区间之中, 并且从这个角度来看, 我们同样无法拒绝原假设。在 10% 的显著性水平下, 我们得到结论: 关于月收益率的总体均值为 1.10% 的命题与我们观测到的 24 个月的数据序列的结果是一致的。注意到 10% 是当假设为真时, 一个相对较高的拒绝对于 1.10% 月度收益率的总体均值假设的概率。

例 7-2 应收账款的下降

时尚设计 (FashionDesigns) 是一家给零售店提供休闲服饰的供应商, 正在关心其客户可能的账款支付下降。管理办公室度量应收账款支付的平均天数的速率。^① 时尚设计 (FashionDesigns) 公司一般保持平均为 45 天的应收账款回收时间。因为经常分析所有公司的应收账款情况花费是很大的, 所以管理办公室利用抽样调查的方法来跟踪客户支付率的情况。一个随机抽取的 50 个账户的样本表明应收账款的平均回收天数为 49 天, 其具有的标准差为 8 天。

(1) 给出与确定证据是否支持所怀疑的关于客户支付速率减缓条件相一致的原假设和备择假设。

(2) 找出第 1 问中对于假设进行检验的检验统计量。

(3) 在 0.05 和 0.01 显著性水平下, 求得第 1 问中假设检验的拒绝点。

(4) 确定是否原假设能在 0.05 和 0.01 显著性水平下被拒绝。

解:

(1) 所怀疑的条件是应收账款回收天数相对于历史上 45 天的回收速率有了增长, 这暗示了一个“大于”的备择假设。 μ 表示应收账款总体平均回收天数, 于是假设为 $H_0: \mu \leq 45$ 对应于 $H_a: \mu > 45$ 。

(2) 因为总体方差是未知的, 我们使用自由度为 $50 - 1 = 49$ 的 t 检验。

(3) 拒绝点在自由度为 49 的一行中去寻找。为了找到在 0.05 显著性水平下的单边拒绝点, 我们利用 0.05 一栏: 其值为 1.677。为了找到在 0.01 显著性水平下的单边拒绝点, 我们利用 0.01 一栏: 其值为 2.405。总之, 在 0.05 显著性水平下, 如果 $t > 1.677$, 我们拒绝原假设; 在 0.01 显著性水平下, 如果 $t > 2.405$, 我们拒绝原假设。

$$(4) \ t_{49} = \frac{49 - 45}{8 / \sqrt{50}} = \frac{4}{1.131\ 371} = 3.536$$

因为 $3.536 > 1.677$, 所以原假设在 0.05 显著性水平下被拒绝。因为 $3.536 > 2.405$, 所以原假设在 0.01 显著性水平下也被拒绝。我们可以很自信地认为时尚设计 (FashionDesigns) 确实面临客户支付放缓的情况。0.01 的显著性水平是一个相对低的拒绝 45 天或者更少假定值的概率。拒绝使得我们对于均值增加到 45 天之上的结论比较自信。

我们上面所述的是当总体方差未知时, 我们对于单个总体均值的检验使用 t 检验。给定至少

① 这个度量表示商家在销售之后, 在获得支付之前必须等待时间的平均长度。该计算方法是 (应收账款)/(每日平均销售额)。

服从近似正态, 当我们使用小样本且我们不知道总体方差时, t 检验是需要的。对于大样本, 中心极限定理表明无论总体的分布是什么, 样本均值近似服从正态分布。因此, t 检验仍旧是合适的, 但是在样本量很大时, 一个替代的检验可能更加有用。

对于大样本, 实务者有时采用 z 检验代替 t 检验来对于一个均值进行检验。[⊖] 在这个情况下使用 z 检验需要有两个条件。首先, 在大样本中, 样本均值应该服从正态分布, 或至少是近似正态分布, 如同我们已经给出的, 以满足 z 检验的正态性假设。其次, t 检验和 z 检验拒绝点之间的差别在样本量较大时, 应该较小。对于在 0.05 显著性水平下的双边检验, z 检验的拒绝点为 1.96 和 -1.96。对于 t 检验, 在 $df=29$ 的拒绝点为 2.045 和 -2.045 (z 检验和 t 检验拒绝点之间的差别约为 4%), 以及在 $df=50$ 下的 2.009 和 -2.009 (z 检验和 t 检验拒绝点之间的差别约为 2.5%)。因为 t 检验已经可以由统计程序输出且对于总体均值未知的情况理论上正确, 所以我们将其作为检验的一种选择。

在非常有限的例子中, 我们可能知道总体的方差; 在这些例子中 z 检验在理论上是正确的。[⊗]

- 替代的 z 检验方法 (the z -test alternative)

1. 如果被抽样的总体服从已知方差 σ^2 的正态分布, 那么对于单个总体均值 μ 的假设检验的检验统计量为

$$z = \frac{\bar{X} - \mu_0}{\sigma / \sqrt{n}} \quad (7-5)$$

2. 如果被抽样的总体的方差未知, 并且样本量较大, 作为 t 检验的替代, 其替代的检验统计量 (根据中心极限定理) 为

$$z = \frac{\bar{X} - \mu_0}{s / \sqrt{n}} \quad (7-6)$$

在上面的等式中, σ 表示已知的总体标准差; s 表示样本标准差; μ_0 表示总体均值的假定值。

当我们使用 z 检验时, 我们最常用来参考的拒绝点列举如下。

- z 检验的拒绝点

A. 在 $\alpha=0.10$ 的显著性水平下。

1. $H_0: \theta = \theta_0$ 对应于 $H_a: \theta \neq \theta_0$ 。拒绝点为 $z_{0.05} = 1.645$ 和 $-z_{0.05} = -1.645$ 。如果 $z > 1.645$ 或者 $z < -1.645$, 就拒绝原假设。

2. $H_0: \theta \leq \theta_0$ 对应于 $H_a: \theta > \theta_0$ 。拒绝点为 $z_{0.10} = 1.28$ 。如果 $z > 1.28$, 就拒绝原假设。

3. $H_0: \theta \geq \theta_0$ 对应于 $H_a: \theta < \theta_0$ 。拒绝点为 $-z_{0.10} = -1.28$ 。如果 $z < -1.28$, 就拒绝原假设。

B. 在 $\alpha=0.05$ 的显著性水平下。

1. $H_0: \theta = \theta_0$ 对应于 $H_a: \theta \neq \theta_0$ 。拒绝点为 $z_{0.025} = 1.96$ 和 $-z_{0.025} = -1.96$ 。如果 $z > 1.96$ 或者 $z < -1.96$, 就拒绝原假设。

⊖ 这些实务者基于样本容量在 t 检验和 z 检验中进行选择。对于小样本 ($n < 30$), 他们使用 t 检验, 而对于大样本, 使用 z 检验。

⊗ 例如, 在蒙特卡罗模拟中, 我们预先设定风险因子的概率分布。如果我们使用一个正态分布, 我们知道真实的均值和方差。蒙特卡罗模拟涉及利用计算机来反映遭受风险的系统的运行; 我们在常见概率分布一章中对于蒙特卡罗模拟进行了讨论。

2. $H_0: \theta \leq \theta_0$ 对应于 $H_a: \theta > \theta_0$ 。拒绝点为 $z_{0.05} = 1.645$ 。如果 $z > 1.645$, 就拒绝原假设。

3. $H_0: \theta \geq \theta_0$ 对应于 $H_a: \theta < \theta_0$ 。拒绝点为 $-z_{0.05} = -1.645$ 。如果 $z < -1.645$, 就拒绝原假设。

C. 在 $\alpha = 0.01$ 的显著性水平下。

1. $H_0: \theta = \theta_0$ 对应于 $H_a: \theta \neq \theta_0$ 。拒绝点为 $z_{0.005} = 2.575$ 和 $-z_{0.005} = -2.575$ 。如果 $z > 2.575$ 或者 $z < -2.575$, 就拒绝原假设。

2. $H_0: \theta \leq \theta_0$ 对应于 $H_a: \theta > \theta_0$ 。拒绝点为 $z_{0.01} = 2.33$ 。如果 $z > 2.33$, 就拒绝原假设。

3. $H_0: \theta \geq \theta_0$ 对应于 $H_a: \theta < \theta_0$ 。拒绝点为 $-z_{0.01} = -2.33$ 。如果 $z < -2.33$, 就拒绝原假设。

例 7-3 发行商业票据对于股票价格的影响

商业票据 (CP) 是无担保的短期公司债券, 如同美国国库券一样, 它是以到期单笔支付为特征的。当一家公司首次进入 CP 市场时, 股票市场的参与者会对所公布的 CP 评级做出何种反映?

纳亚尔 (Nayar) 和罗瑟夫 (Rozeff) (1994) 利用 1981 年 10 月到 1985 年 12 月区间的数据回答了这个问题。在这个区间中, 有 132 个 CP 发行者 (96 个工业企业以及 36 个非工业企业) 接受了标准普尔信用周刊或者穆迪投资者服务债券调查的初始评级。纳亚尔 (Nayar) 和罗瑟夫 (Rozeff) 将评级归为高信用或低信用。高信用 CP 评级为标准普尔的 A1+ 或者 A1 和穆迪的 Prime-1 (P1)。低信用的 CP 评级为标准普尔的 A2 或者更低评级以及穆迪的 Prime-2 (P2) 或者更低评级。初始评级的公布日定义为 $t=0$ 。然而, 研究者发现公司本身经常在信用周刊或者债券调查公布评级之前, 散布评级信息。对于股价反映的研究从公布日前 $t-1$ 开始, 因为这天接近于实际信息发布的日期。

如果 CP 评级提供了对于股价估值的新的有用信息, 那么信息应该在其获得的同时, 就造成股价和收益率的一个变动。只有股票收益率这一个因素是我们感兴趣的: 超出给定市场风险或者贝塔的收益率, 被称为超常收益率。正的 (负的) 超常收益率表明投资者在评级公告中感受到了对于公司有利的信息。虽然纳亚尔 (Nayar) 和罗瑟夫 (Rozeff) 检验了不同时间区间或时间窗口的超常收益率, 我们这里仅摘选了他们在评级公告之前日期 ($t-1$) 研究的一部分结果进行报告:

所有 CP 发行 ($n=132$ 次发行)。原假设为在 $t-1$ 这天的平均超常收益率为 0。如果股票投资者在公告中没有发现正面的或者负面的信息, 那么原假设将为真。

$$\text{平均超常收益率} = 0.39\%$$

$$\text{平均超常收益率的样本标准误} = 0.1336\% \ominus$$

工业企业的 CP 发行具有高信用评级 ($n=72$ 次发行)。原假设为在 $t-1$ 这天的平均超常股票收益率为 0。如果股票投资者在公告中没有发现正面的或者负面的信息, 那么原假设将为真。

$$\text{平均超常收益率} = 0.79\%$$

$$\text{平均超常收益率的样本标准误} = 0.197\%$$

工业企业的 CP 发行具有低信用评级 ($n=24$ 次发行)。原假设为在 $t-1$ 这天的平均超常股票收益率为 0。如果股票投资者在公告中没有发现正面的或者负面的信息, 那么原假设将为真。

⊖ 这个标准误是由计算 132 个发行的样本标准差 (一个截面的样本标准差) 除以 132 的平方根得到的。其他的标准误也是通过同样方法得到的。

平均超常收益率 = -0.57%

平均超常收益率的样本标准误 = 0.38%

研究者选择使用 z 检验。

(1) 对于这 3 种情况中的每一种，如果假设原假设反映了对于投资者在初始评级中普遍地没有发现正面的或者负面的信息的一个信念。给出一组假设（一个原假设和一个备择假设）包含所有这 3 种情况。

(2) 对于所有 CP 发行的情况，确定在第 1 问中所给出的原假设在 0.05 和 0.01 显著性水平下是否应该被拒绝。解释你所得到的结果。

(3) 对于工业企业具有高信用评级的 CP 发行情况，确定在第 1 问中所给出的原假设在 0.05 和 0.01 显著性水平下是否应该被拒绝。解释你所得到的结果。

(4) 对于工业企业具有低信用评级的 CP 发行情况，确定在第 1 问中所给出的原假设在 0.05 和 0.01 显著性水平下是否应该被拒绝。解释你所得到的结果。

解：

(1) 与股票投资者相关的 CP 信用评级中没有信息一致的一组假设为

H_0 ：在 $t-1$ 时点的总体超常平均收益率等于 0

H_a ：在 $t-1$ 时点的总体超常平均收益率不等于 0

(2) 从 z 检验拒绝点的信息来看，我们了解到如果 $z > 1.96$ 或者如果 $z < -1.96$ ，在 0.05 显著性水平下我们就应该拒绝原假设，而如果 $z > 2.575$ 或者如果 $z < -2.575$ ，在 0.01 显著性水平下我们就应该拒绝原假设。利用 z 检验， $z = (0.39\% - 0\%) / 0.1336\% = 2.92$ ，其在 0.05 和 0.01 显著性水平下都是显著的。故原假设被拒绝。CP 发行这个事实本身被认为是利好消息。

因为显著的结果可能是由于异常点所造成的，所以研究者们也报告了正超常收益率和负超常收益率的观测数。正、负超常收益率观测数之比为 80:52，这个结果倾向于支持 z 检验中得到的出现正超常收益率的结论。

(3) 利用 z 检验， $z = (0.79\% - 0\%) / 0.197\% = 4.01$ ，其在 0.05 和 0.01 显著性水平下都是显著的。对于高信用的初始 CP 评级这一消息，股票表现出明显正的超常收益率。投资者可能将评级机构所给予高的评级而确信公司未来的良好前景。

正、负超常收益率的观测数之比为 48:24，这个结果倾向于支持 z 检验中得到的出现正超常收益率的结论。

(4) 利用 z 检验， $z = (-0.57\% - 0\%) / 0.38\% = -0.15$ ，其无论在 0.05 显著性水平下还是在 0.01 显著性水平下都不显著。在低信用评级的情況下，我们无法得到投资者在初始 CP 评级公告之后发现或者是正面的或者是负面的信息的结论。

正、负超常收益率的观测数之比为 11:13，这个结果倾向于支持 z 检验中得到的结论，即无法拒绝原假设。

几乎所有的实际情况都涉及未知总体方差的情况。表 7-2 汇总了在总体方差未知情况下，我们对于总体均值检验的讨论。

表 7-2 对于总体均值的检验（总体方差未知）

	大样本（ $n \geq 30$ ）	小样本（ $n < 30$ ）
总体服从正态分布	t 检验（可选择使用 z 检验）	t 检验
总体不服从正态分布	t 检验（可选择使用 z 检验）	未知

7.3.2 对均值间差异的检验

我经常想知道两个组的均值（例如，一个平均收益率）之间是否是不同的。我们所观测到的不同是否由于碰巧，还是因为它们的均值确实存在不同？我们具有两个样本，每个样本都来自一个组类。当我们有理由认为样本是来自服从至少是近似正态分布的总体，并且样本间是相互独立的时候，那么我们就可以应用本节中所介绍的方法。我们讨论检验两个总体均值间是否不同的两种 t 检验。在一种情况下，总体方差虽然是未知的，但是可以被假设为相等的。那么，我们将来自两个样本的观测值有效地放在一起，得到一个合并的但总体方差未知的共同估计值。一个合并（pooled）估计量是从两个不同样本的合并样本中所得到的估计量。在第2种情况下，我们不能假设未知的总体方差是相等的，于是我们获得了一个近似的 t 检验。分别令代表第1个总体的均值 μ_1 和第2个总体的均值 μ_2 不变，我们最经常想要检验的是总体均值间是否相等或者一个是否比另一个大。因此，我们经常给出下列的假设：

- 1. $H_0: \mu_1 - \mu_2 = 0$ 对应于 $H_a: \mu_1 - \mu_2 \neq 0$ （备择假设为 $\mu_1 \neq \mu_2$ ）
- 2. $H_0: \mu_1 - \mu_2 \leq 0$ 对应于 $H_a: \mu_1 - \mu_2 > 0$ （备择假设为 $\mu_1 > \mu_2$ ）
- 3. $H_0: \mu_1 - \mu_2 \geq 0$ 对应于 $H_a: \mu_1 - \mu_2 < 0$ （备择假设为 $\mu_1 < \mu_2$ ）

然而，我们可以给出其他的假设，诸如 $H_0: \mu_1 - \mu_2 = 2$ 对应于 $H_a: \mu_1 - \mu_2 \neq 2$ 。但是检验的方法是相同的。

t 检验的定义如下：

- 关于两个总体均值间差别检验的检验统计量（test statistic for a test of the difference between two population means）（正态分布总体，总体方差未知但是假设相等（normally distributed populations, population variances unknown but assumed equal））。当我们能假设两个总体服从正态分布并且未知的总体方差相等时，基于独立随机样本的 t 检验给出如下：

$$t = \frac{(\bar{X}_1 - \bar{X}_2) - (\mu_1 - \mu_2)}{\left(\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}\right)^{1/2}} \tag{7-7}$$

式中， $s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$ 是合并共同方差的估计值。自由度的数值为 $n_1 + n_2 - 2$ 。

例 7-4 标准普尔 500 的平均收益率：一个跨度为 10 年的等式检验

20 世纪 80 年代的标准普尔 500 指数已实现的月度平均收益率似乎与 20 世纪 70 年代的月度平均收益率有着非常巨大的不同。那么该不同是否在统计上是显著的呢？表 7-3 中所给出的数据表明，我们没有充足的理由拒绝这两个 10 年的收益率的总体方差是相同的。

表 7-3 两个 10 年的标准普尔 500 指数的月度平均收益率及其标准差

10 年区间	月份数 (n)	月度平均收益率 (%)	标准差
20 世纪 70 年代	120	0.580	4.598
20 世纪 80 年代	120	1.470	4.738

- (1) 给出与双边假设检验相一致的原假设和备择假设。
- (2) 找出检验第 1 问中假设的检验统计量。
- (3) 求出第 1 问中所检验的假设在 0.10、0.05 和 0.01 显著性水平下的拒绝点。
- (4) 确定在 0.10、0.05 和 0.01 显著性水平下是否应该拒绝原假设。

解:

(1) 令 μ_1 表示 20 世纪 70 年代的总体平均收益率, 而 μ_2 表示 80 年代的总体平均收益率, 于是我们给出如下的假设:

$$H_0: \mu_1 - \mu_2 = 0 \text{ 对应于 } H_a: \mu_1 - \mu_2 \neq 0$$

(2) 因为两个样本分别取自不同的 10 年区间, 所以它们是独立样本。总体方差是未知的, 但是可以被假设为相等。给定所有这些条件, 在式 (7-7) 中所给出的 t 检验具有 $120 + 120 - 2 = 238$ 的自由度。

(3) 在数值表 (附录 B) 中, 最接近 238 的自由度为 200。对于一个双边检验, $df = 200$ 的 0.10、0.05 和 0.01 显著性水平下的拒绝点分别为 ± 1.653 、 ± 1.972 和 ± 2.601 。总而言之, 在 0.10 显著性水平下, 如果 $t < -1.653$ 或者 $t > 1.653$, 我们将拒绝原假设; 在 0.05 显著性水平下, 如果 $t < -1.972$ 或者 $t > 1.972$, 我们将拒绝原假设; 而在 0.01 显著性水平下, 如果 $t < -2.601$ 或者 $t > 2.601$, 我们将拒绝原假设。

(4) 在计算检验统计量时, 第 1 步是计算合并方差的估计值:

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2} = \frac{(120 - 1)(4.598)^2 + (120 - 1)(4.738)^2}{120 + 120 - 2}$$

$$= \frac{5\,187.239\,512}{238} = 21.795\,124$$

$$t = \frac{(\bar{X}_1 - \bar{X}_2) - (\mu_1 - \mu_2)}{\left(\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}\right)^{1/2}} = \frac{(0.580 - 1.470) - 0}{\left(\frac{21.795\,124}{120} + \frac{21.795\,124}{120}\right)^{1/2}} = \frac{-0.89}{0.602\,704} = -1.477$$

t 值 -1.477 在 0.10 显著性水平下不显著, 因此它同样在 0.05 和 0.01 显著性水平下也不显著。因此, 我们无法在任一个显著性水平下拒绝原假设。

在许多我们所关注的实际例子中, 我们无法假设总体方差是相等的。下面的检验量经常被用于这一情况下的投资文献中。

- 关于两个总体均值间差别检验的检验统计量 (test statistic for a test of the difference between two population means) (正态分布总体, 总体方差未知且不相等 (normally distributed populations, unequal and unknown population variances))。当我们能假设两个总体服从正态分布, 但是不知道总体方差, 而且不能假设方差是相等的时候, 基于独立随机样本的近似 t 检验给出如下:

$$t = \frac{(\bar{X}_1 - \bar{X}_2) - (\mu_1 - \mu_2)}{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^{1/2}} \quad (7-8)$$

其中, 我们使用“修正的”自由度, 其计算公式为

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{(s_1^2/n_1)^2}{n_1} + \frac{(s_2^2/n_2)^2}{n_2}} \quad (7-9)$$

的数值表。

一个应用的诀窍是在计算自由度之前先计算 t 检验量。 t 检验量是否是显著的, 有时会很

明显。

例 7-5 违约债券的回收率：一个假设检验

具有风险的公司债券的要求收益率是如何来决定的？两个重要的考虑因素为预期违约概率和在违约发生的情况下预期能够回收的金额，或者说回收率。奥尔特曼（Altman）和基肖尔（Kishore）（1996）首次记录了以行业和信用等级进行分层的违约债券的平均回收率。对于他们的研究区间 1971 ~ 1995 年，奥尔特曼（Altman）和基肖尔（Kishore）发现公共事业公司、化工类公司、石油公司以及塑胶制造公司的违约债券的回收率明显要高于其他行业。这一差别是否能够通过在高回收率行业中的高信用债券比较来解释？他们通过检验以信用等级分层的回收率来对此进行研究。这里，我们仅讨论他们对于高信用担保债券的结果。其中 μ_1 表示公共事业公司的高信用担保债券的总体平均回收率，而 μ_2 表示其他行业（非公共事业）公司的高信用担保债券的总体平均回收率，假设为 $H_0: \mu_1 - \mu_2 = 0$ 对应于 $H_a: \mu_1 - \mu_2 \neq 0$ 。

表 7-4 摘自他们的部分结果。

表 7-4 高信用债券的回收率

行业类/高信用	公共事业样本			非公共事业样本		
	观测数	平均价格 ^①	标准差	观测数	平均价格 ^①	标准差
公共事业 高信用担保	21	64.42 美元	14.03 美元	64	55.75 美元	25.17 美元

资料来源：奥尔特曼（Altman）和基肖尔（Kishore）（1996），表 5。

①这是在违约时的平均价格，是对于回收率的一个度量指标。

根据研究者们的假设，总体（公共事业公司的回收率和非公共事业公司的回收率）服从正态分布，并且样本间是独立的。基于表中的数据，回答下列问题：

- (1) 讨论为什么奥尔特曼和基肖尔会选择基于式 (7-8)，而不是式 (7-7) 的检验方法。
- (2) 计算检验上述给出的原假设的检验统计量。
- (3) 该检验的修正自由度的数值为多少？
- (4) 确定在 0.10 显著性水平下是否应该拒绝原假设。

解：

(1) 高信用担保的公共事业公司回收率的样本标准差（14.03 美元）要比与之相比较的非公共事业公司回收率的样本标准差（25.17 美元）更小。故不假设它们的均值相等的选择是恰当的，所以奥尔特曼和基肖尔采用式 (7-8) 所给出的近似 t 检验。

(2) 检验统计量为

$$t = \frac{(\bar{X}_1 - \bar{X}_2)}{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^{1/2}}$$

式中， $\bar{X}_1$ 表示公共事业公司的样本平均回收率 = 64.42； $\bar{X}_2$ 表示非公共事业公司的样本平均回收率 = 55.75； s_1^2 表示公共事业公司的样本方差 = $14.03^2 = 196.8409$ ； s_2^2 表示非公共事业公司的样本方差 = $25.17^2 = 633.5289$ ； n_1 表示公共事业公司的样本大小 = 21； n_2 表示非公共事业公司的样本大小 = 64。

因此, $t = (64.42 - 55.75) / [(196.8409/21) + (633.5289/64)]^{1/2} = 8.67 / (9.373376 + 9.898889)^{1/2} = 8.67 / 4.390019 = 1.975$ 。所以所计算得到的 t 统计量为 1.975。

(3)

$$\begin{aligned} df &= \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{(s_1^2/n_1)^2}{n_1} + \frac{(s_2^2/n_2)^2}{n_2}} = \frac{\left(\frac{196.8409}{21} + \frac{633.5289}{64}\right)^2}{\frac{(196.8409/21)^2}{21} + \frac{(633.5289/64)^2}{64}} \\ &= \frac{371.420208}{5.714881} = 64.99 \text{ 或者 } 65 \text{ 个自由度} \end{aligned}$$

(4) 在 t 分布的数值表中最接近 $df = 65$ 的一栏是 $df = 60$ 。对于 $\alpha = 0.10$, 我们找到 $t_{\alpha/2} = 1.671$ 。因此, 如果 $t < -1.671$ 或者 $t > 1.671$, 我们会拒绝原假设。基于所算出的值 1.975, 我们在 0.10 显著性水平下拒绝原假设。存在一些公共事业公司和非公共事业公司回收率不同的证据。为什么是这样的? 奥尔特曼和基肖尔认为公司资产的不同性质以及不同行业的竞争水平造成了不同的回收率情况。

7.3.3 对(配对样本)均值差的检验

在之前的一节中, 我们给出了辨别总体均值间差别的两种 t 检验。这两个检验是基于两个样本的。一个对于那两个检验有效性的假设为两个样本间是独立的, 即它们之间没有关联。当我们想要对于我们认为是相关的样本间的两个均值进行检验时, 就要应用本节中所介绍的方法。

本节中的 t 检验是基于配对观测值 (paired observations) 形式的数据, 所以这个检验本身有时也被称为配对比较检验 (paired comparisons test)。配对观测值是相关的观测值, 因为它们具有某些共同特点。配对比较检验是一个对于相关样本间差别的检验。例如, 我们可能关心在税收法律变化之前和之后的公司红利政策对于红利税收的影响。于是, 同一家公司具有“之前”和“之后”的一对观测值。我们可以检验关于我们所观测的同一家公司的前后均值间差别的假设。在其他情况下, 配对观测值不是针对同一个观测值对象的。例如, 我们可能检验两个投资策略所获得的平均收益率是否在所研究区间内相等。这里的观测值之间是相关的, 这是基于在每个月对于每个策略有一个观测值, 而这两个观测值之间依赖于给定的市场风险因素。因为两个策略的收益率可能与某个共同的风险因素相关, 例如市场收益率, 所以这两个样本间是相关的。通过计算基于均值间差别的标准误, 我们可以给出如下考虑观测值相关性的 t 检验。

令 A 表示“之后”, B 表示“之前”, 假设我们的观测值为随机变量 X_A 和 X_B , 并且这两个样本是相关的。我们将观测值以配对方式进行组织。令 d_i 表示两个配对观测值之间的差别。我们可以使用记号 $d_i = x_{Ai} - x_{Bi}$, 其中, x_{Ai} 和 x_{Bi} 是对于第 i 对两个随机变量的观测值, $i = 1, 2, \dots, n$ 。令 μ_d 表示总体均值的差别。我们可以给出如下假设, 其中, μ_{d0} 是总体均值差别的假定值:

1. $H_0: \mu_d = \mu_{d0}$ 对应于 $H_a: \mu_d \neq \mu_{d0}$
2. $H_0: \mu_d \leq \mu_{d0}$ 对应于 $H_a: \mu_d > \mu_{d0}$
3. $H_0: \mu_d \geq \mu_{d0}$ 对应于 $H_a: \mu_d < \mu_{d0}$

在实际应用中, μ_{d0} 最常用的值为 0。

通常, 我们考虑方差未知的正态总体, 并且我们将给出一个 t 检验。为了计算 t 统计量, 我

们首先要求得样本平均差别:

$$\bar{d} = \frac{1}{n} \sum_{i=1}^n d_i \tag{7-10}$$

式中, n 表示配对观测值的数目。样本方差, 记为 s_d^2 , 为

$$s_d^2 = \frac{\sum_{i=1}^n (d_i - \bar{d})^2}{n - 1} \tag{7-11}$$

对于这个量取平方根, 我们就有样本的标准差 s_d , 其使得我们能够如下计算平均差别的标准误:[⊖]

$$s_{\bar{d}} = \frac{s_d}{\sqrt{n}} \tag{7-12}$$

- 对于平均差别检验的检验统计量 (test statistic for a test of mean differences) (正态分布总体, 未知总体方差 (normally distributed populations, unknown population variances))。当我们具有取自方差未知正态总体所产生的样本的配对观测值所组成的数据时, t 检验是基于

$$t = \frac{\bar{d} - \mu_{d0}}{s_{\bar{d}}} \tag{7-13}$$

其具有 $n-1$ 自由度, 式中, n 表示配对观测值的数目, $\bar{d}$ 表示样本平均差别 (由式 (7-10) 所给出的), $s_{\bar{d}}$ 表示 $\bar{d}$ 的标准误 (由式 (7-12) 所给出的)。

表 7-5 报告了两个对于稀有金属进行专门投资的投资组合在 1997 ~ 2002 年的季度收益率情况。两个投资组合在风险上 (以收益率的标准差和其他指标进行度量) 密切相似, 而且几乎具有相同的支出比率。一家著名的投资服务公司在 2003 年年初对投资组合 B 的评级设定要高于投资组合 A。在调查投资组合相对业绩表现时, 假设我们想要检验在 1997 ~ 2002 年投资组合 A 的平均季度收益率等于投资组合 B 的平均季度收益率的假设。因为两个投资组合本质上具有相同的一组风险因素, 所以它们的收益率不是相互独立的, 因此, 使用配对比较检验是恰当的。令 μ_d 表示这个区间中两个投资组合收益率间差别的总体平均值。我们在 0.05 显著性水平下检验 $H_0: \mu_d = 0$ 对应于 $H_a: \mu_d \neq 0$ 。

表 7-5 两个管理的投资组合的季度收益率 (1997 ~ 2002 年)

季度	投资组合 A (%)	投资组合 B (%)	差别 (投资组合 A-投资组合 B)
4Q: 2002	11.40	14.64	-3.24
3Q: 2002	-2.17	0.44	-2.61
2Q: 2002	10.72	19.51	-8.79
1Q: 2002	38.91	50.40	-11.49
4Q: 2001	4.36	1.01	3.35
3Q: 2001	5.13	10.18	-5.05
2Q: 2001	26.36	17.77	8.59
1Q: 2001	-5.53	4.76	-10.29
4Q: 2000	5.27	-5.36	10.63
3Q: 2000	-7.82	-1.54	-6.28

⊖ 我们也可以使用下面等价的表达式, 其利用了两个变量间的相关系数: $s_d^2 = \sqrt{s_A^2 + s_B^2 - 2r(X_A, X_B) s_A s_B}$, 其中 s_A^2 是 X_A 的样本方差, s_B^2 是 X_B 的样本方差, 而 $r(X_A, X_B)$ 是 X_A 和 X_B 之间的样本相关系数。

(续)

季度	投资组合 A (%)	投资组合 B (%)	差别 (投资组合 A-投资组合 B)
2Q: 2000	2.34	0.19	2.15
1Q: 2000	-14.38	-12.07	-2.31
4Q: 1999	-9.80	-9.98	0.18
3Q: 1999	19.03	26.18	-7.15
2Q: 1999	4.11	-2.39	6.50
1Q: 1999	-4.12	-2.51	-1.61
4Q: 1998	-0.53	-11.32	10.79
3Q: 1998	5.06	0.46	4.60
2Q: 1998	-14.01	-11.56	-2.45
1Q: 1998	12.50	3.52	8.98
4Q: 1997	-29.05	-22.45	-6.60
3Q: 1997	3.60	0.10	3.50
2Q: 1997	-7.97	-8.96	0.99
1Q: 1997	-8.62	-0.66	-7.96
均值	1.87	2.52	-0.65

差别的样本标准差 = 6.71

投资组合 A 和投资组合 B 之间的样本差别的平均值 $\bar{d}$ 为每季度 -0.65%。样本差别平均值的标准误为 $s_{\bar{d}} = \frac{6.71}{\sqrt{24}} = 1.369\,673$ 。所计算出的检验统计量 $t = (-0.65 - 0) / 1.369\,673 = -0.475$ ，其具有 $n - 1 = 24 - 1 = 23$ 的自由度。在 0.05 的显著性水平下，如果 $t > 2.069$ 或者 $t < -2.069$ ，我们就拒绝原假设。因为 -0.475 小于 -2.069，故我们无法拒绝原假设。在 0.10 显著性水平下，如果 $t > 1.714$ 或者 $t < -1.714$ ，我们就拒绝原假设。因此，季度平均收益率的差别在任何常用显著性水平下都不显著。

下面的例子说明了对于两个竞争的投资策略进行评估的这个检验的应用。

■ 例 7-6 道-10 投资策略

麦奎因 (McQueen)、谢尔德斯 (Shields) 和索利 (Thorley) (1997) 检验了一个流行的投资策略 (该策略投资于道琼斯工业平均指数中收益率最高的 10 只股票) 与一个买入并持有的策略 (该策略投资于道琼斯工业平均指数中所有的 30 只股票) 之间的业绩比较。他们研究的区间段是 1946 ~ 1995 年的 50 年区间。

从表 7-6 中，我们有 $\bar{d} = 3.06\%$ 和 $s_{\bar{d}} = 6.62\%$ 。

表 7-6 道-10 和道-30 投资组合年度收益率汇总 (1946 ~ 1995 年) ($n = 50$)

策略	平均收益率	标准差
道-10	16.77%	19.10%
道-30	13.71	16.64
差别	3.06	6.62 ^①

①差别的样本标准差。
资料来源：麦奎因，谢尔德斯和索利 (1997)，表 1。

(1) 给出与道-10 和道-30 策略间收益率差别的均值等于 0 这个双边检验相一致的原假设和备择假设。

(2) 找出对于第 1 问中假设进行检验的检验统计量。

(3) 求出在 0.01 显著性水平下第 1 问中所检验的假设的拒绝点。

(4) 确定在 0.01 显著性水平下是否应该拒绝原假设。(使用本书后所给出的数值表。)

(5) 讨论为什么选择配对比较检验。

解:

(1) μ_d 表示道-10 和道-30 策略间收益率差别的均值, 我们有 $H_0: \mu_d = 0$ 对应于 $H_a: \mu_d \neq 0$ 。

(2) 因为总体方差未知, 所以检验统计量为一个自由度为 $50 - 1 = 49$ 的 t 检验。

(3) 在 t 分布的数值表中, 我们查阅自由度为 49 的一行, 显著性水平为 0.05 的一列, 从而得到 2.68。如果我们发现 $t > 2.68$ 或者 $t < -2.68$, 我们将拒绝原假设。

$$(4) t = \frac{3.06}{6.62/\sqrt{50}} = \frac{3.06}{0.936209} = 3.2685 \text{ 或者 } 3.27$$

因为 $3.27 > 2.68$, 所以我们拒绝原假设。作者得到结论: 平均收益率的差别在统计上是明显显著的。然而, 在经过对于道-10 的高风险、额外交易成本和不利的税收政策的调整之后, 他们发现道-10 投资组合并没有在经济意义上击败道-30。

(5) 道-30 包含道-10。因此, 它们不是相互独立的样本; 通常, 道-10 和道-30 间收益率的相关系数为正。因为样本是相互依赖的, 配对比较检验是恰当的。

7.4 关于方差的假设检验

因为方差和标准差被广泛地用于投资中对于风险的度量, 所以分析师们应该熟悉对于方差的假设检验。本节讨论的检验在投资文献中经常出现。我们将考察两种类型的检验: 关于单个总体方差值的检验和关于两个总体方差间差别的检验。

7.4.1 对单个方差的检验

在这一节中, 我们讨论对于单个总体的方差值 σ^2 的假设检验。我们使用 σ_0^2 来表示 σ^2 的假定值。我们给出的假设如下:

1. $H_0: \sigma^2 = \sigma_0^2$ 对应于 $H_a: \sigma^2 \neq \sigma_0^2$ (一个“不等于”的备择假设)
2. $H_0: \sigma^2 \leq \sigma_0^2$ 对应于 $H_a: \sigma^2 > \sigma_0^2$ (一个“大于”的备择假设)
3. $H_0: \sigma^2 \geq \sigma_0^2$ 对应于 $H_a: \sigma^2 < \sigma_0^2$ (一个“小于”的备择假设)

在关于单个正态分布总体方差的检验中, 我们使用卡方检验, 记为 χ^2 。 χ^2 分布不同于正态分布和 t 分布, 其是非对称的。与 t 分布一样, χ^2 分布是一族分布。对于每个自由度的可能值 $n-1$ (n 是样本容量) 都存在一个不同的分布。但不同于 t 分布, χ^2 分布以 0 为下界; χ^2 不会取到负值。

- 对于总体方差值检验的检验统计量 (test statistic for tests concerning the value of a population variance) (正态总体 (normal population))。如果我们具有来自一个正态分布总体的 n 个独立的观测值, 恰当的检验量为

$$\chi^2 = \frac{(n-1)s^2}{\sigma_0^2} \quad (7-14)$$

其具有 $n-1$ 的自由度。在表达式的分子中的是样本方差, 其计算为

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1} \quad (7-15)$$

与 t 检验不同的是, 例如, χ^2 检验对于违背其假设的情况非常敏感。如果样本实际上并不是随机的, 或者如果样本并不来自于一个正态分布总体, 那么基于 χ^2 检验的推断将很可能是错误的。

如果我们选择显著性水平 α , 那么这 3 类假设的拒绝点给出如下:

- 总体方差假设检验的拒绝点。

1. “不等于” H_a : 如果检验统计量大于自由度为 $df=n-1$ 的 χ^2 分布上侧 $\alpha/2$ 点 (记为 $\chi^2_{\alpha/2}$) 或者小于其下侧 $\alpha/2$ 点 (记为 $\chi^2_{1-\alpha/2}$), 那么我们会拒绝原假设。[⊖]

2. “大于” H_a : 如果检验统计量大于自由度为 $df=n-1$ 的 χ^2 分布上侧 α 点, 那么我们就拒绝原假设。

3. “小于” H_a : 如果检验统计量小于自由度为 $df=n-1$ 的 χ^2 分布下侧 α 点, 那么我们就拒绝原假设。

例 7-7 一家股票型共同基金的风险收益特征 (2)

你继续对于森达 (Sendar) 股票型基金的分析, 该基金是一家成立只有 24 个月的中市值成长型基金。回忆一下, 在这个区间中, 森达股票型基金实现了 3.60% 的月度收益率的样本标准差。你现在想要检验根据森达特定的投资策略是否会导致一个小于 4% 的月度收益率的标准差。

- (1) 给出与研究目标一致的原假设和备择假设。

- (2) 找出对于第 1 问中假设进行检验的检验统计量。

- (3) 求出在 0.05 显著性水平下对于第 1 问所检验假设的拒绝点。

- (4) 确定在 0.05 显著性水平下是否应该拒绝原假设。(使用本书书后所给出的数值表。)

解:

- (1) 我们有一个“小于”的备择假设, 其中, σ 是森达股票型基金的相关标准差。注意要将标准差平方, 以得到检验中所表达的方差, 该假设为 $H_0: \sigma^2 \geq 16.0$ 对应于 $H_a: \sigma^2 < 16.0$ 。

- (2) 检验统计量为自由度为 $24-1=23$ 的 χ^2 。

- (3) 下侧 0.05 的拒绝点在 $df=23$ 的一行, 0.95 的一列 (95% 的右侧尾部概率, 给出了检验统计量大于等于该拒绝点的概率为 0.95) 被找到。该拒绝点为 13.091。如果 χ^2 小于 13.091, 那么我们将拒绝原假设。

$$(4) \chi^2 = \frac{(n-1) s^2}{\sigma_0^2} = \frac{23 \times 3.60^2}{4^2} = \frac{298.08}{16} = 18.63$$

因为 18.63 (所计算的检验统计量的值) 不小于 13.091, 所以我们无法拒绝原假设。我们不能得出森达的投资策略导致小于 4% 的月度收益率标准差的结论。

⊖ 正如其他假设检验, χ^2 检验可以给出一个置信区间的解释。然而, 不同于基于 z 统计量和 t 统计量的置信区间, 对于方差的 χ^2 置信区间是不对称的。基于样本量为 n 的对于总体方差的一个双边置信区间, 具有下界 $L = (n-1) s^2 / \chi^2_{\alpha/2}$ 和上界 $U = (n-1) s^2 / \chi^2_{1-\alpha/2}$ 。在原假设下, 总体方差的假定值应该落入这两个边界之中。

7.4.2 对两个方差是否相等的检验

假设我们有一个关于两个均值分别为 μ_1 和 μ_2 , 方差为 σ_1^2 和 σ_2^2 的正态分布总体的方差相对值的假设。我们可以给出如下所有假设选择中的一个:

1. $H_0: \sigma_1^2 = \sigma_2^2$ 对应于 $H_a: \sigma_1^2 \neq \sigma_2^2$
2. $H_0: \sigma_1^2 \leq \sigma_2^2$ 对应于 $H_a: \sigma_1^2 > \sigma_2^2$
3. $H_0: \sigma_1^2 \geq \sigma_2^2$ 对应于 $H_a: \sigma_1^2 < \sigma_2^2$

注意, 在等式的这种情况下, 原假设 $\sigma_1^2 = \sigma_2^2$ 暗示了总体方差的比率等于 1: $\sigma_1^2/\sigma_2^2 = 1$ 。给定从这些总体中所得到的独立随机样本, 关于这些假设的检验是基于 F 检验, 其是一个样本方差的比率。假设我们使用 n_1 个观测值来计算样本方差 s_1^2 而用 n_2 个观测值来计算样本方差 s_2^2 。那么关于两个总体方差间差别的检验将使用 F 分布。与 χ^2 分布相同, F 分布也是一族具有下界为 0 的不对称的分布。每个 F 分布是由两个自由度的值所定义的, 它们分别称为分子和分母的自由度。[⊖] F 检验与 χ^2 检验一样, 对于其假设违背的情况是不稳健的。

- 对于两个总体均值间差别检验的检验统计量 (test statistic for tests concerning differences between the variances of two populations) (正态分布总体 (normally distributed populations))。假设我们有两个样本, 第 1 个样本有 n_1 个观测值并且其样本方差为 s_1^2 , 第 2 个样本有 n_2 个观测值并且其样本方差为 s_2^2 。样本是随机的且相互之间独立, 而且都是由正态分布总体所产生的。基于样本方差比率的两个总体方差间差别的检验为

$$F = \frac{s_1^2}{s_2^2} \quad (7-16)$$

其具有 $df_1 = n_1 - 1$ 的分子自由度, $df_2 = n_2 - 1$ 的分母自由度。注意 df_1 和 df_2 是用来分别计算 s_1^2 和 s_2^2 的除数。

传统上或者常用的处理手法是使用 $\frac{s_1^2}{s_2^2}$ 或者 $\frac{s_2^2}{s_1^2}$ 这两个比率中较大者作为实际检验统计量。当我们遵照这一传统处理方法时, 检验统计量的值总是大于等于 1; 于是, F 临界值的数值表只需要包含大于等于 1 的值。在该传统处理方法下, 任一给出假设的拒绝点是相关 F 分布右侧的单个数值。注意总体的编号“1”或“2”在任一例子中是任意的。

- 两个总体方差相对值的假设检验的拒绝点。根据使用 $\frac{s_1^2}{s_2^2}$ 或者 $\frac{s_2^2}{s_1^2}$ 这两个比率中较大值的传统处理方法, 考虑下面两种情况:

1. 一个“不等于”的备择假设: 如果检验统计量大于某个设定的分子和分母自由度的 F 分布的上侧 $\alpha/2$ 点, 那么我们就在 α 显著性水平下拒绝原假设。
2. 一个“大于”或“小于”的备择假设: 如果检验统计量大于某个设定的分子和分母自由度的 F 分布的上侧 α 点, 那么我们就在 α 显著性水平下拒绝原假设。

[⊖] χ^2 分布和 F 分布之间的关系如下: 如果 χ_1^2 是一个自由度为 m 的 χ^2 随机变量, χ_2^2 是一个自由度为 n 的 χ^2 随机变量, 那么 $F = (\chi_1^2/m)/(\chi_2^2/n)$ 就服从分子自由度为 m 、分母自由度为 n 的 F 分布。

因此，如果我们在 $\alpha = 0.01$ 显著性水平下进行一个双边检验，我们需要在 F 数值表中寻找在单边检验下 $\alpha/2 = 0.01/2 = 0.005$ 显著性水平下的拒绝点（第 1 种情况）。但是在 0.01 显著性水平下的一个单边检验使用 F 数值表中 $\alpha = 0.01$ 的拒绝点（第 2 种情况）。举一个例子，假设我们要在 0.05 显著性水平下进行一个双边检验。利用书后 $\alpha/2 = 0.05/2 = 0.025$ 的 F 数值表，我们发现其拒绝点为 2.72。因为 2.77 这个值大于 2.72，我们在 0.05 显著性水平下拒绝原假设。

如果不遵照上面所规定的传统处理方法，并且我们给出一个小于 1 的 F 统计量的计算值，那么我们是否仍旧会使用 F 数值表？这个回答是肯定的；利用 F 统计量的倒数性质，我们可以计算所需的值。最简单给出该性质的方法是通过计算来说明。假设我们对于一个双边检验选择 0.05 的显著性水平，而且我们的 F 值为 0.11，其分子自由度为 7，分母自由度为 9。我们取其倒数， $1/0.11 = 9.09$ 。于是我们查阅对于 0.025 显著性水平（因为这是一个双边检验）相反自由度的 F 数值表中的该值：对于分子自由度为 9，分母自由度为 7 的 F 值。换句话说， $F_{9,7} = F_{7,9}$ ，并且 9.09 超过了 4.82 的临界值，所以 $F_{7,9} = 0.11$ 在 0.05 显著性水平下是显著的。

例 7-8 波动率和 1987 年股市暴跌

你正在调查 1987 年 10 月市场暴跌之后标准普尔 500 收益率的总体方差是否发生了改变。你收集了表 7-7 中给出的 1987 年 10 月之前的 120 个月和 1987 年 10 月之后的 120 个月的收益率数据。你将显著性水平设定为 0.01。

表 7-7 在 1987 年 10 月之前与之后的标准普尔 500 收益率及其方差

	n	月度平均收益率 (%)	收益率方差
1987 年 10 月之前	120	1.498	18.776
1987 年 10 月之后	120	1.392	13.097

- (1) 给出与研究目的的语言描述相一致的原假设和备择假设。
- (2) 找出对于第 1 问中假设进行检验的检验统计量。
- (3) 确定在 0.01 显著性水平下是否应该拒绝原假设。（使用本书书后的 F 数值表。）

解：

- (1) 我们有一个“不等于”的备择假设：

$$H_0: \sigma^2_{\text{之前}} = \sigma^2_{\text{之后}} \text{ 对应于 } H_a: \sigma^2_{\text{之前}} \neq \sigma^2_{\text{之后}}$$

- (2) 为了检验两个方差相等的原假设，我们使用 $F = \frac{s_1^2}{s_2^2}$ ，其分子和分母自由度均为 $120 - 1 = 119$ 。

(3) “之前”的样本方差更大，所以根据计算 F 统计量的传统方法，“之前”的样本方差出现在分子中： $F = 18.776/13.097 = 1.434$ 。因为这是一个双边检验，我们使用 0.005 (0.01/2) 显著性水平下 F 数值表中的值来给出了一个 0.01 显著性水平下的检验值。在书后的数值表中，最接近 119 自由度的值是 120 自由度。在 0.01 显著性水平下，拒绝点是 1.61。因为 1.434 小于 1.61 这个临界值，所以我们无法拒绝原假设，即无法认为收益率的总体方差在暴跌之前与之后的两个区间是相同的。

例 7-9 衍生品到期日的波动率

在 20 世纪 80 年代的美国，人们开始对于 3 个衍生品（股票期权，指数期权和期货）在一年的 4 个月中，同时在一到期的情况予以了关注。这样的日子被称为“三重魔力日”。表 7-8

给出了在 1983 年 7 月 1 日至 1986 年 10 月 24 日区间段中正常日和期权/期货到期日的日度标准差数据。表中的数据是关于标的为标准普尔 100（标准普尔 500 的一个子集，其中包括 100 个流动性最好的在标准普尔 500 中所交易的具有可交易期权的股票）的期权和期货。

表 7-8 收益率的标准差：1983 年 7 月 1 日 ~ 1986 年 10 月 24 日

日期的类型	<i>n</i>	标准差 (%)
正常交易日	115	0.786
期权/期货到期日	12	1.178

资料来源：基于爱德华兹 (Edwards) (1988) 中的表 1。

- (1) 给出与三重魔力日表现出超过正常交易日波动率的认识相一致的原假设和备择假设。
- (2) 给出用来对于第 1 问中的假设进行检验的检验统计量。
- (3) 确定在 0.05 的显著性水平下是否应该拒绝原假设。(利用书后的 *F* 分布数值表。)

解：

- (1) 我们具有一个“大于”的备择假设：

$$H_0: \sigma^2_{\text{到期日}} \leq \sigma^2_{\text{正常交易日}} \text{ 对应于 } H_a: \sigma^2_{\text{到期日}} > \sigma^2_{\text{正常交易日}}$$

(2) 根据之前给出的对于分子和分母选择的传统，令 σ_1^2 表示三重魔力日的方差，而 σ_2^2 表示正常交易日的方差。为了检验原假设，我们使用 $F = s_1^2/s_2^2$ ，其分子的自由度为 $12 - 1 = 11$ ，分母的自由度为 $115 - 1 = 114$ 。

(3) $F = \frac{(1.178)^2}{(0.786)^2} = \frac{1.388}{0.618} = 2.25$ 。因为这是一个在 0.05 显著性水平下的单侧检验，我们

直接使用 *F* 分布表中的 0.05 显著性水平。在书后的表中，最接近 114 自由度的值是 120。在 0.05 显著性水平下，拒绝域的临界点为 1.87。因为 2.25 比 1.87 要大，所以我们拒绝原假设。它表明三重魔力日的波动率要高于正常交易日。

7.5 其他议题：非参数推断

至此，我们所讨论的假设检验具有两个共同的特征。首先，它们都是关于参数的；其次，它们的有效性取决于一组给定的假设。例如，均值和方差是两个正态分布的参数或者定义的量。检验也做出具体的假设，特别是关于产生样本总体分布的假设。任何具有上面两个特征之一的检验或者方法是一个参数检验 (parametric test) 或者参数方法。然而，在一些例子中，我们关注的是除分布参数以外的量。而在另一些例子中，我们可能认为参数检验的假设对于我们所使用的数据并不满足。在这种情况下，非参数检验或者方法可能会非常有用。一个非参数检验 (nonparametric test) 是指与参数无关的或者是对于样本所来自的总体做出最少假设的一种检验方法。

我们主要在 3 种情况下使用非参数方法：当我们使用的数据不符合分布假设时，当我们的数据是以排序给出时，或者当假设中我们没有考虑参数时。

当分析师所获得的数据表明参数的分布假设没有得到满足时，就会出现第 1 种情况。例如，我们可能想要检验一个关于总体均值的假设，但是我们认为无论是 *t* 检验还是 *z* 检验都不适合，因为样本很小，而且可能来自明显的非正态分布总体。在这个例子中，我们可以使用非参数检验。非参数检验经常涉及将观测值（或者观测值的一个函数）根据数据大小转化为排序，而且有时它只涉及“大于”或者“小于”的相互关系（利用“+”和“-”符号来表示这些关系）。

典型地，我们必须参考专门的统计数值表来确定检验统计量的拒绝临界值，至少对于小样本来说是这样的。[Ⓐ]于是，这类检验通常是将原假设表示为排序或者正负号的一个命题。在表 7-9 中，我们给出了不同于我们在本章之前所讨论的参数检验的一个非参数的例子。[Ⓑ]读者应该参考一本介绍这些检验的综合性商业统计教科书或者一本专业的教科书来了解更加具体的内容。[Ⓒ]

表 7-9 替代参数检验的非参数检验

	参数检验	非参数检验
关于单个均值的检验	t 检验 z 检验	威尔科克森 (Wilcoxon) 符号秩检验
关于均值之间差别的检验	t 检验 近似 t 检验	曼恩 - 惠特尼 (Mann-Whitney) U 检验
关于均值差的检验 (配对比较检验)	t 检验	威尔科克森 (Wilcoxon) 符号秩检验 符号检验

当使用非参数检验时，我们经常将原始数据转化为排序。在某些情况下，原始数据本来就是排序。在这种情况下，我们也使用非参数检验，因为参数检验一般获得一个比排序带有更强的测量尺度的结果。例如，如果我们的数据是投资经理的排序，关于这些排序的假设可以利用非参数方法来进行检验。排序数据也出现在许多其他的金融环境中。例如、希尼 (Heaney)、古贺 (Koga)、奥利弗 (Oliver) 和郑 (Tran) (1999) 研究日本公司的规模 (以收入来度量) 与它们对于衍生品的使用之间的关系。被研究的公司使用衍生品来对冲一个或者 5 种风险暴露中的多个：利率风险，外汇风险，商品价格风险，可交易证券价格风险和信用风险。研究者们对于每家公司给出一个“风险暴露的可感知范围”的评分，该评分等于该公司所报告的面临风险暴露的种数。虽然收入是以一个具体的尺度来度量的 (一个比率尺度)，而风险暴露的范围是以顺序尺度来度量的。[Ⓓ]因此，研究者使用非参数统计量来研究衍生品的使用和公司规模之间的关系。

我们会使用非参数检验方法的第 3 种情况是在我们的问题不是关于参数时发生。例如，如果我们的问题是关于一个样本是否随机的，那么我们使用一个适当的非参数检验 (即所谓的游程检验)。另一类非参数的问题表述为一个样本是否来自一个服从某个特定概率分布的总体 (例如，利用科尔莫戈罗夫 - 斯米尔诺夫 (Kolmogorov-Smirnov) 检验)。

在本章的最后，我们通过稍稍具体地介绍一个经常在投资研究中使用的非参数统计量——斯皮尔曼秩相关系数来结束这一章的内容。

7.5.1 相关性检验：斯皮尔曼 (Spearman) 秩相关系数

在许多投资的背景下，我们想要评估两个变量间线性关系的强度，即两者间的相关系数。在大部分的情况下，我们使用在概率论中的一些概念一章和相关性和回归一章中所讨论的相关系数。然而，基于相关系数的关于两个变量不相关假设的 t 检验依赖于相当严格的假设。[Ⓔ]当我们认

Ⓐ 对于大样本，经常存在检验统计量的一个转化，使得其能够使用标准正态分布或者 t 分布的数值表。
Ⓑ 在某些情况下，对于一个参数检验存在多个不同的非参数方法。
Ⓒ 例如参见 Hettmansperger 和 McKean (1998) 或者 Siegel (1956)。
Ⓓ 我们在统计学概念和市场收益率一章中对于度量的尺度进行了讨论。
Ⓔ t 检验在相关性和回归一章中进行讨论。检验的假设为关于两个变量 (X, Y) 的每个观测值 (x, y) 是来自二元正态分布的一个随机观测值。非正式地，在一个二元或者两个变量的正态分布中，每个变量服从正态分布，并且它们的联合关系完全由它们之间的相关系数 ρ 所决定的。想了解更多内容，参见 Daniel 和 Terrell (1986)。

为所考虑的总体偏离这些假设时，我们可以使用一个基于斯皮尔曼秩相关系数（Spearman rank correlation coefficient） r_s 的检验。斯皮尔曼秩相关系数本质上等于通常的由各个样本内两个变量的秩所计算得到的相关系数。因此，它是一个介于 -1 到 +1 之间的量，其中 -1(+1) 分别表示变量间的完全负（正）的线性相关，而 0 表示没有任何线性关系（不相关）。 r_s 的计算需要以下步骤：

- 1. 将基于 X 的观测值从大到小进行排序。分配数字 1 给具有最大数值的观测值，将数字 2 给予第 2 大的观测值，以此类推。如果有多个相同的值，我们将每个相同的观测值给予它们所共同占据的排序的平均值。例如，如果第 3 大和第 4 大的值是相同的，那么我们将给予这两个观测值 3.5 的排序（3 和 4 的平均值）。对于基于 Y 的观测值也进行相同的处理。
- 2. 计算每对 X 和 Y 观测值间的差别 d_i 。
- 3. 于是，在样本容量为 n 的情况下，斯皮尔曼秩相关系数为：[Ⓐ]

$$r_s = 1 - \frac{6 \sum_{i=1}^n d_i^2}{n(n^2 - 1)}$$

(7-17)

假设一个投资者想要投资于一个美国大市值的成长型共同基金。他已经将选择的范围缩小到 10 家基金之中。在检验这些基金时，关于基金所报告的 3 年的夏普比率是否与它们最近报告的支出比率相关的问题出现在了该投资者的面前。因为基于相关系数 t 检验的假设可能并不满足，所以进行一个基于秩相关系数的检验是恰当的。[Ⓑ]表 7-10 给出了 r_s 的计算过程。[Ⓒ]

表 7-10 斯皮尔曼至相关系数：一个例子

	共同基金									
	1	2	3	4	5	6	7	8	9	10
夏普比率 (X)	-1.08	-0.96	-1.13	-1.16	-0.91	-1.08	-1.18	-1.00	-1.06	-1.00
支出比率 (Y)	1.34	0.92	1.02	1.45	1.35	0.50	1.00	1.50	1.45	1.50
按 X 的排名	6.5	2	8	9	1	6.5	10	3.5	5	3.5
按 Y 的排名	6	9	7	3.5	5	10	8	1.5	3.5	1.5
d_i	0.5	-7	1	5.5	-4	-3.5	2	2	1.5	2
d_i^2	0.25	49	1	30.25	16	12.25	4	4	2.25	4
$r_s = 1 - \frac{6 \sum d_i^2}{n(n^2 - 1)} = 1 - \frac{6(123)}{10(100 - 1)} = 0.2545$										

表中的前面两行包含了原始的数据。按 X 排名的一行将夏普比率转化为了排名；而按 Y 排名的一行将支出比率转化为了排名。我们想要检验 $H_0: \rho = 0$ 对应于 $H_a: \rho \neq 0$ ，其中 ρ 在这里被定义为在 X 和 Y 经过排序之后的总体的相关系数。对于小样本来说，基于 r_s 的该检验的拒绝临界值必须查阅表 7-11。而对于大样本来说（例如， $n > 30$ ），我们利用下式进行 t 检验：

Ⓐ 基于排序的通常的相关系数的计算会得到与式 (7-17) 近似相同的结果。

Ⓑ 支出比率（基金运营支出与平均净资产间的比率）同时具有下界（在 0 以下）和上界。在实务中，夏普比率的观测值也是出现在一个有限的区间之中。因此，没有一个变量是服从正态分布的，因此他们不可能共同服从二元正态分布。简而言之， t 检验的假设无法得到满足。

Ⓒ 表中的数据是基于标准普尔对于实际大市值成长型基金至 2003 年第一季度为止的 3 年区间中的共同基金报告中的统计值所给出的。负的夏普比率反映了在这个区间中部分下降的美国股票市场情况。

$$t = \frac{(n - 2)^{1/2} r_s}{(1 - r_s^2)^{1/2}}$$

(7-18)

基于 $n - 2$ 个自由度。

表 7-11 斯皮尔曼秩相关系数近似分布的上侧拒绝域临界点

样本大小: n	$\alpha=0.05$	$\alpha=0.025$	$\alpha=0.01$	样本大小: n	$\alpha=0.05$	$\alpha=0.025$	$\alpha=0.01$
5	0.800 0	0.900 0	0.900 0	18	0.399 4	0.471 6	0.548 0
6	0.771 4	0.828 6	0.885 7	19	0.389 5	0.457 9	0.533 3
7	0.678 6	0.745 0	0.857 1	20	0.378 9	0.445 1	0.520 3
8	0.619 0	0.714 3	0.809 5	21	0.368 8	0.435 1	0.507 8
9	0.583 3	0.683 3	0.766 7	22	0.359 7	0.424 1	0.496 3
10	0.551 5	0.6364	0.733 3	23	0.351 8	0.415 0	0.485 2
11	0.527 3	0.609 1	0.700 0	24	0.343 5	0.406 1	0.474 8
12	0.496 5	0.580 4	0.671 3	25	0.336 2	0.397 7	0.465 4
13	0.478 0	0.554 9	0.642 9	26	0.329 9	0.389 4	0.456 4
14	0.459 3	0.534 1	0.622 0	27	0.323 6	0.382 2	0.448 1
15	0.442 9	0.517 9	0.600 0	28	0.317 5	0.374 9	0.440 1
16	0.426 5	0.500 0	0.582 4	29	0.311 3	0.368 5	0.432 0
17	0.411 8	0.485 3	0.563 7	30	0.305 9	0.362 0	0.425 1

注：相对应的下侧拒绝域的临界值是通过将上侧拒绝域临界的符号改变而得到的。

在目前的这个例子中，具有 0.05 显著性水平的双侧检验，由表 7-11 给出了 $n = 10$ 时的上侧拒绝域临界值为 0.636 4（我们使用 0.025 一行作为在 0.05 显著性水平下双侧检验的值）。相应地，如果 r_s 小于 $-0.636 4$ 或者大于 0.636 4，我们就应该拒绝原假设。由于 r_s 等于 0.254 5，所以我们无法拒绝原假设。

在共同基金的例子中，我们将两个观测值变量转化为排序值。如果一个或者两个原始数据已经是排序值的形式给出的话，我们将需要使用 r_s 来检验其相关性。

7.5.2 非参数推断：总结

非参数统计的方法扩展了统计推断所能应用的领域，因为非参数的方法没有做出什么假设，并且可以应用于排序数据以及回答与参数不相关的问题。在很多情况下，非参数检验是和参数检验结果一起进行报告的。由此，读者就能够了解所报告的统计结果对于参数检验所做出假设的敏感性到底有多大。然而，如果参数检验的假设实际上是满足的话，（能获得的）参数检验一般要优于非参数检验，因为参数检验通常能使我们得到更加精确的结论。[⊖] 如果要了解所有在金融和投资相关文献中所可能遇到的非参数检验方法的全貌，最好的方式是去查阅一本专业的教科书。[⊖]

⊖ 使用先前章节中所介绍的概念，参数检验一般来说更加有用。
⊖ 例如，可以参见 Hettmansperger 和 McKean（1998）或者 Siegel（1956）。

相关性和回归

8.1 引言

作为一名金融分析师，你经常需要检验两个或者两个以上金融变量之间的关系。例如，你可能想要了解不同股票市场指数间的收益率是否相关，如果是相关的，它们又是以何种方式相关的。或者你可能假设一家公司投资资本收益率与其资本成本之间的差额能够帮助我们解释该公司在市场中的价值。相关性与回归分析是检验这些问题的工具。

本章组织安排如下：在第 2 节中，我们将介绍相关性分析，这是一个度量两个变量是如何相互变化的工具。我们讨论的问题包括相关系数的计算、解释、应用、局限性以及统计检验。第 3 节介绍回归分析中的基本概念，一个非常有用的用来检验一个或者多个变量（自变量）解释或者预测另一个变量（因变量）能力的方法。

8.2 相关性分析

我们有许多方法来检验两组数据是如何相关的。两个最有用的方法是散点图和相关性分析。我们首先讨论散点图。

8.2.1 散点图

散点图（scatter plot）是一个在二维坐标中显示两组数据序列观测值之间关系的图形。例如，假设我们想要画出 6 个工业化国家的长期货币增长和长期通货膨胀之间的关系，以考察这两个变量间相关性有多强。表 8-1 给出了 1970 ~ 2001 年间 6 个国家的货币供应量平均年度增长率和平均年度通货膨胀率。

为了将表 8-1 中的数据转化为一个散点图，我们将每个国家的数据在图上以一个点标出。对于每个点，其 x 轴坐标表示的是 1970 ~ 2001 年年度平均货币供应量增长率，而 y 轴坐标表示的是 1970 ~ 2001 年年度平均通货膨胀率。图 8-1 给出了由表 8-1 的数据所绘制的散点图。

表 8-1 各国年度货币供应量增长率和年度通货膨胀率 (1970 ~ 2001 年)

国家	货币供应量增长率 (%)	通货膨胀率 (%)
澳大利亚	11.66	6.76
加拿大	9.15	5.19
新西兰	10.60	8.15
瑞士	5.75	3.39
英国	12.58	7.58
美国	6.34	5.09
平均	9.35	6.03

资料来源：国际货币基金组织。

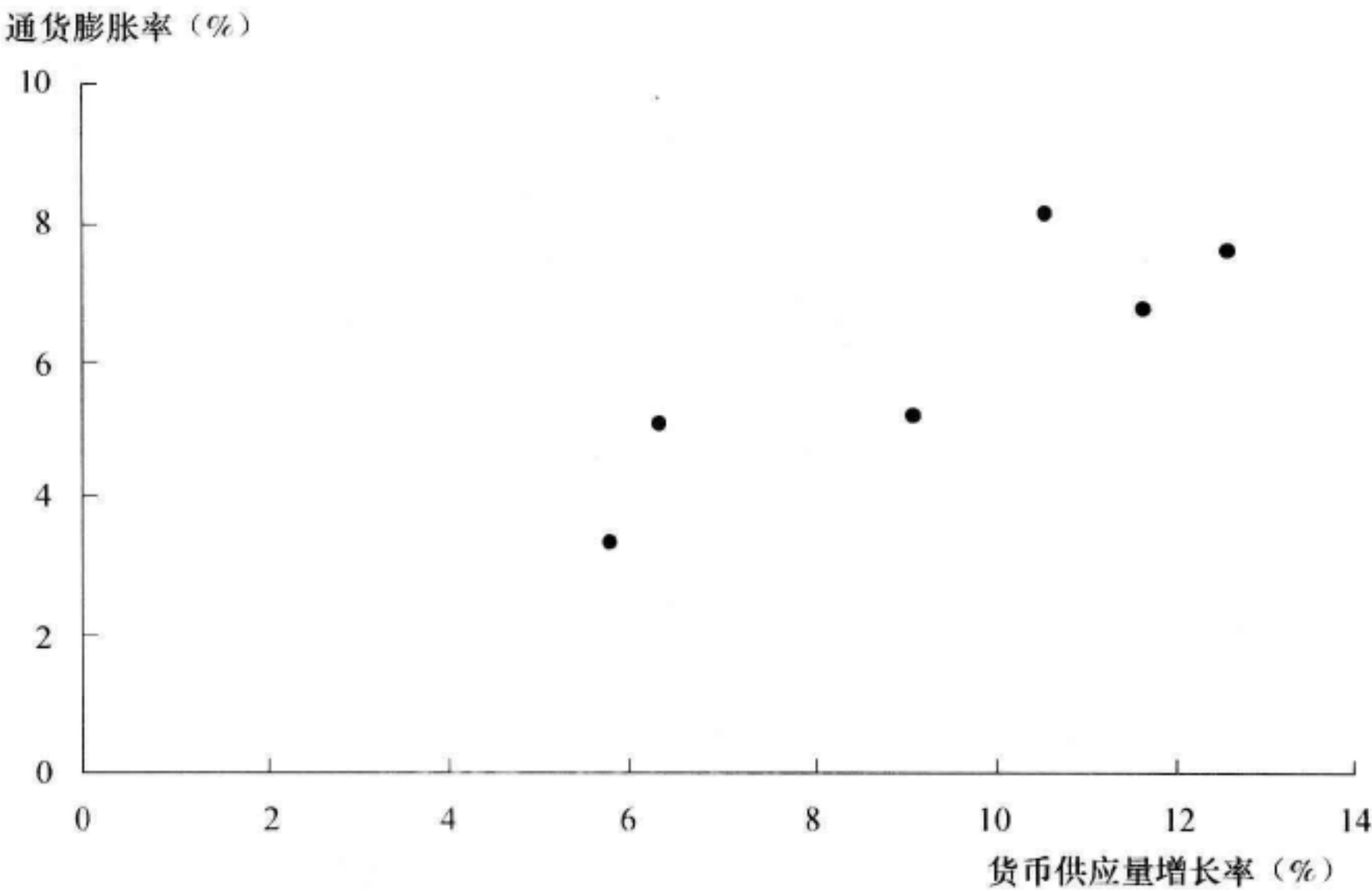


图 8-1 各国年度货币供应量增长率和通货膨胀率的散点图 (1970 ~ 2001 年)

资料来源：国际货币基金组织。

注意散点图中的每个观测值都是以一个点来表示的，这些点之间是不相连的。散点图并没有显示出哪个观测值是来自哪个国家；它只在图上显示出两组数据序列配对的实际观测值。例如，图中最右边的点表示的是代表英国的数据。由数据所绘制的图 8-1 表现出相当强的具有正斜率的线性关系。下面我们讨论如何将这一线性关系进行定量。

8.2.2 相关性分析

与散点图（用图像描绘两个数据序列间的关系）不同，相关性分析（correlation analysis）会利用单个数来反映这一相同的关系。相关系数是一个度量两组数据序列相关程度有多密切的指标。特别地，相关系数度量了两个变量间线性相关（linear association）的方向和程度。一个相关系数可能具有的最大值为 1，最小值为 -1。一个大于 0 的相关系数表示两个变量间的一个正的线性关系：当一个变量增加（或减少）时，另一个变量倾向于增加（或减少）。一个小于 0 的相关系数表示两个变量间的一个负的线性关系：当一个变量增加（或减少）时，另一个变量倾向于减少（或增加）。一个等于 0 的相关系数表示两个变量间不存在线性关系。[⊖]图 8-2 显示两个相关系数为 1 的变量的散点图。

⊖ 后面内容中，我们会显示具有相关系数为 0 的两个变量可能具有较强的非线性关系。

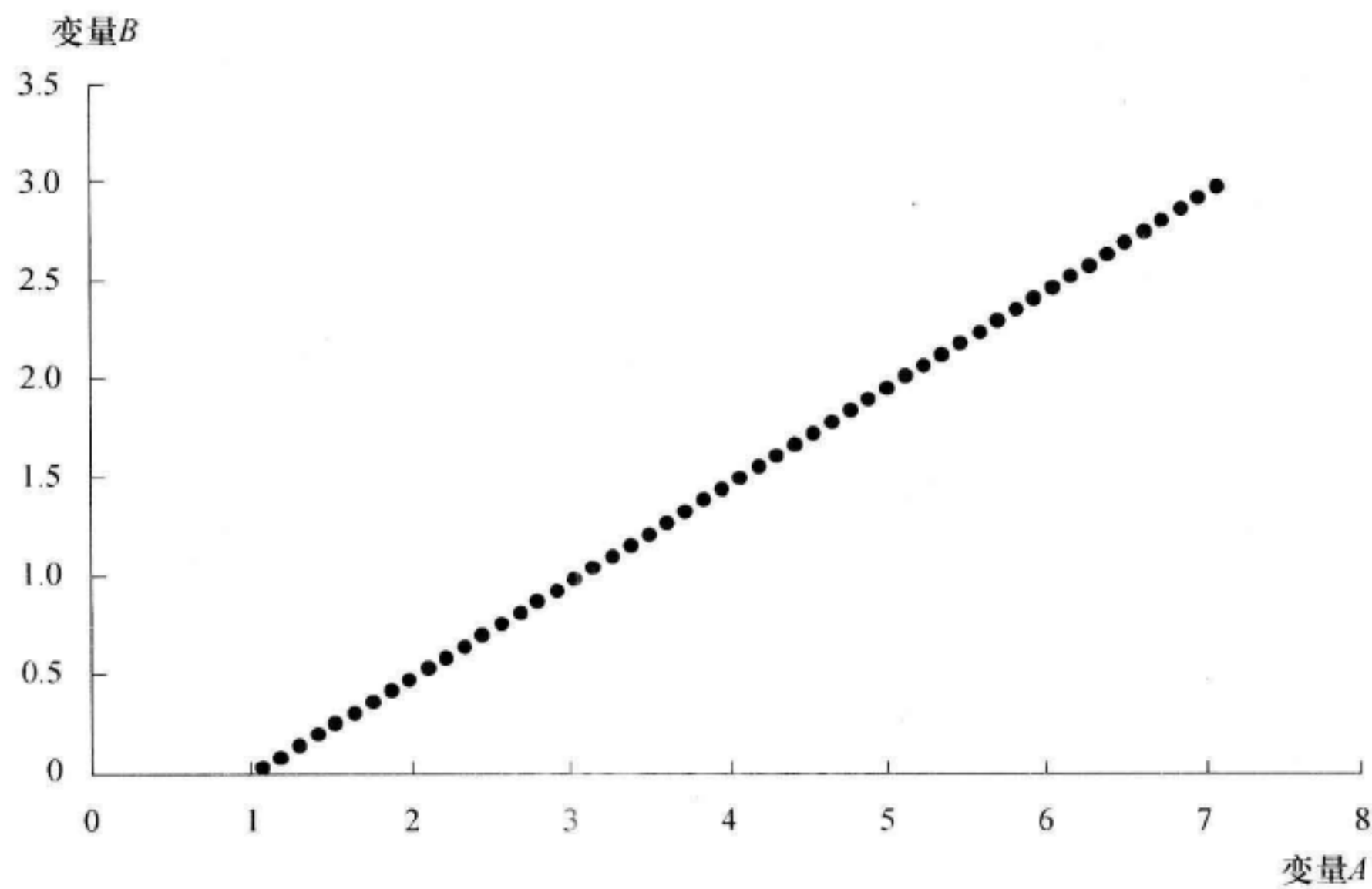


图 8-2 相关系数为 1 的两个变量

注意在图 8-2 中散点图上的所有点都在一条正的斜率的直线上。当变量 A 增长一个单位，变量 B 将增加半个单位。因为图上的所有点都在一条直线上， A 增加一单位将与 B 增加相同的半个单位精确相关，而无论 A 处于什么水平。即使图中该直线的斜率是不同的（但是为正），只要所有的点都在一条直线上，那么这两个变量的相关性将为 1。

图 8-3 表明两个相关系数为 -1 的变量的散点图。同样地，绘制的每个观测点落在一条直线上。然而，在这个图中，该直线具有负的斜率。随着 A 增长一单位， B 将减少半个单位，而无论最初 A 的值为多少。

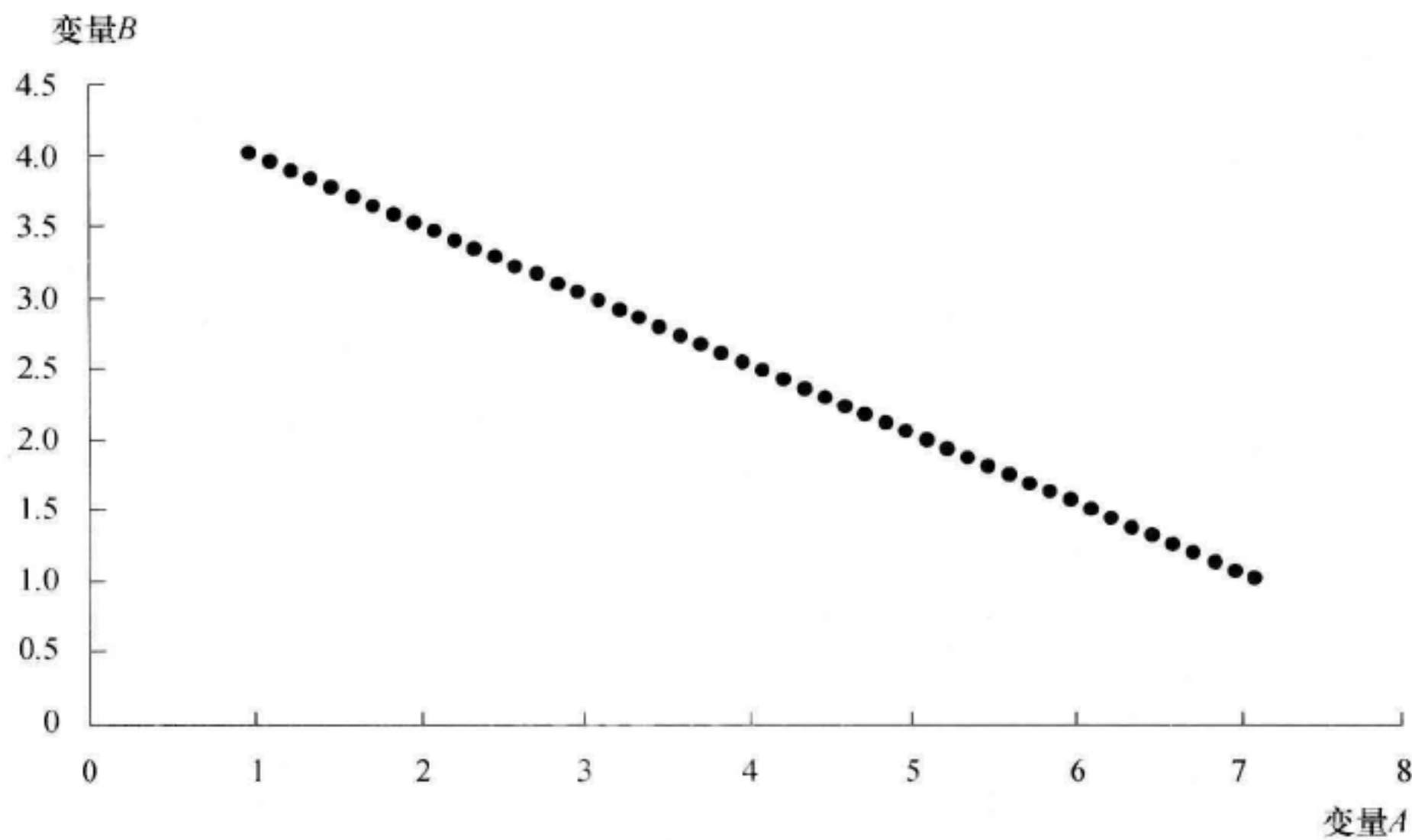


图 8-3 相关系数为 -1 的两个变量

图 8-4 表明两个相关系数为 0 的变量的散点图；它们没有线性关系。该图像表明 A 的值并没有告诉我们任何与 B 的值有关的信息。

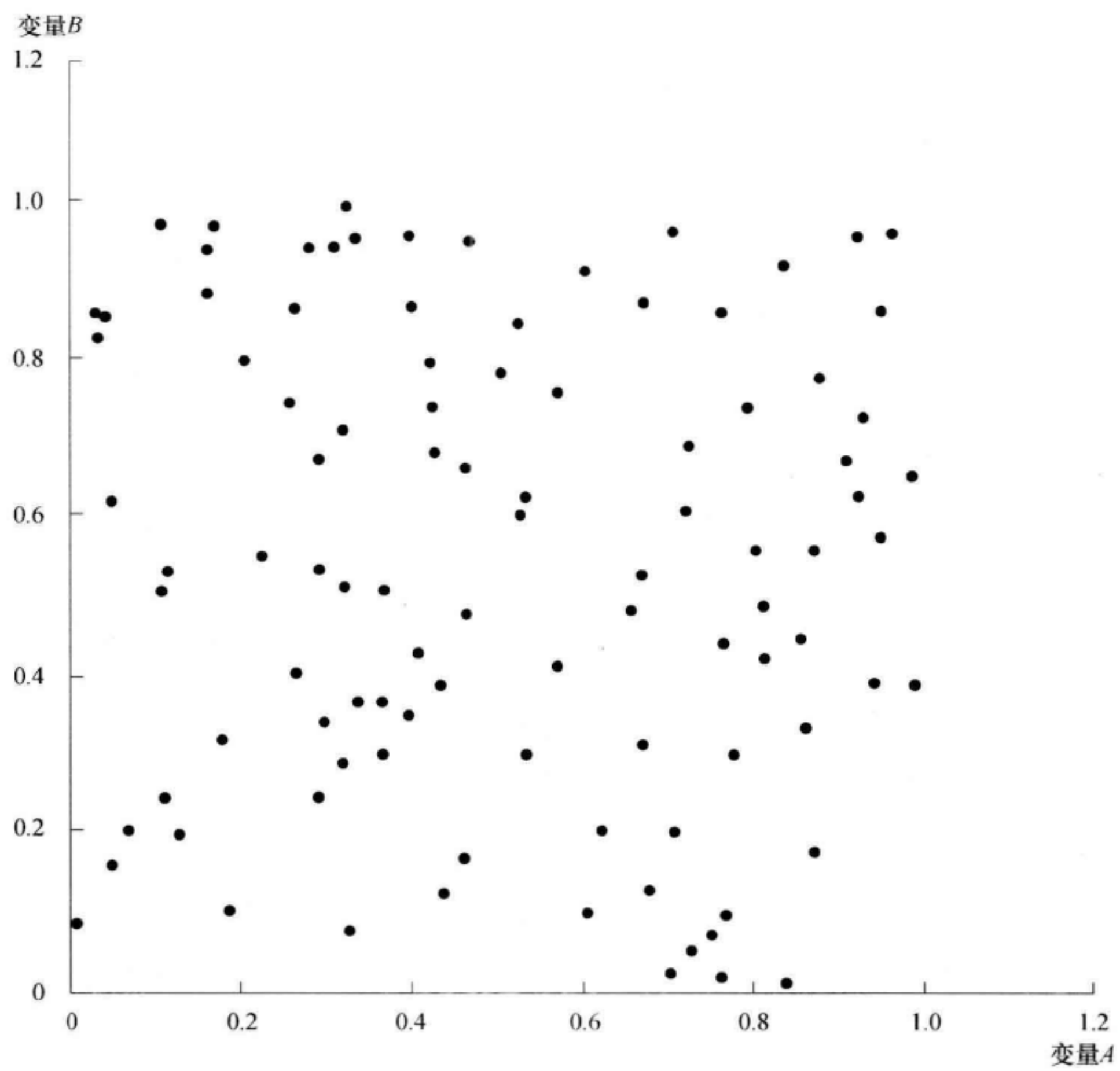


图 8-4 相关系数为 0 的两个变量

8.2.3 计算和解释相关系数

为了定义和计算相关系数，我们需要另一个度量线性关系的量：协方差。在概率论中的一些概念一章中，我们将协方差定义为两个随机变量偏离它们各自总体均值偏差乘积的期望值。这是总体协方差的定义，我们将其也用在对于未来的预测中。为了研究历史的或者样本的协方差，我们需要使用样本协方差。对于一个样本量为 n 的样本， X 和 Y 的样本协方差为

$$\text{Cov}(X,Y) = \sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})/(n-1) \tag{8-1}$$

样本协方差是两个随机变量的观测值偏离它们样本均值偏差值乘积的平均值。[⊖]如果随机变量是收益率，那么协方差的单位将是收益率的平方。

样本相关系数比样本协方差要容易解释。为了理解样本相关系数，我们需要随机变量 X 的样本标准差的表达式。我们需要计算 X 的样本方差来得到其样本标准差。一个随机变量的方差可以简单地视为该随机变量与自己的协方差。 X 的样本方差 s_X^2 的表达式为

$$s_X^2 = \sum_{i=1}^n (X_i - \bar{X})^2/(n-1)$$

⊖ 在分母中使用 $n-1$ 是一个技术要点，它保证了样本协方差是一个总体协方差的无偏估计量。

样本标准差是样本方差正的平方根：

$$s_X = \sqrt{s_X^2}$$

样本方差和样本标准差都是度量观测值对于样本均值离散度的指标。标准差使用与随机变量相同的单位；而方差是以其单位的平方来度量的。

计算样本相关系数的公式为

$$r = \frac{\text{Cov}(X,Y)}{s_X s_Y} \tag{8-2}$$

相关系数是两个变量（ X 和 Y ）的协方差除以它们样本标准差（ s_X 和 s_Y ）的乘积。与协方差一样，相关系数是对于线性关系的一个度量指标。然而，相关系数具有的优势在于它是一个没有度量单位的纯数。因为相关系数是通过将协方差除以标准差的乘积得到的，所以它没有单位。因为我们在本章中将使用样本方差、样本标准差和协方差，所以我们将对这些统计量进行多次计算。

表 8-2 显示了如何从表 8-1 中的数据计算组成相关系数等式（式（8-2））中的不同成分。[⊖]对于 1970 ~ 2001 年各国年度平均货币供应量增长率的单个观测值被记为 X_i ，而对于 1970 ~ 2001 年各国年度平均通货膨胀率的单个观测值被记为 Y_i 。余下的列显示了计算相关系数的输入变量：样本协方差和样本标准差。

表 8-2 样本协方差和样本标准差：各国年度货币供应量增长率和年度通货膨胀率（1970 ~ 2001 年）

国家	货币供应量增长率 X_i	通货膨胀率 Y_i	交叉乘积 $(X_i - \bar{X})(Y_i - \bar{Y})$	偏差的平方 $(X_i - \bar{X})^2$	偏差的平方 $(Y_i - \bar{Y})^2$
澳大利亚	0.116 6	0.067 6	0.000 169	0.000 534	0.000 053
加拿大	0.091 5	0.051 9	0.000 017	0.000 004	0.000 071
新西兰	0.106 0	0.081 5	0.000 265	0.000 156	0.000 449
瑞士	0.057 5	0.033 9	0.000 950	0.001 296	0.000 697
英国	0.125 8	0.075 8	0.000 501	0.001 043	0.000 240
美国	0.063 4	0.050 9	0.000 283	0.000 906	0.000 088
总和	0.560 8	0.361 6	0.002 185	0.003 939	0.001 598
平均	0.093 5	0.060 3			
协方差			0.000 437		
方差				0.000 788	0.000 320
标准差				0.028 071	0.017 889

注：1. 将交叉乘积总和除以 $n-1$ ($n=6$) 得到 X 和 Y 的协方差。
2. 将偏差平方总和除以 $n-1$ ($n=6$) 得到 X 和 Y 的方差。
资料来源：国际货币基金组织。

利用表 8-2 中给出的数据，我们能计算出这两个变量的样本相关系数，具体计算如下：

$$r = \frac{\text{Cov}(X,Y)}{s_X s_Y} = \frac{0.000\,437}{(0.028\,071)(0.017\,889)} = 0.870\,2$$

相关系数近似为 0.87，这表明样本中各国长期货币供应量增长率和长期通货膨胀率之间有一个较强的线性关系。相关系数以数值方式反映了这个较强的相关关系，而图 8-1 中的散点图以图形的形式给出了该信息。

计算相关系数所必需的假设是什么呢？如果 X 和 Y 的均值和方差以及 X 和 Y 的协方差都是有

⊖ 我们没有使用表中计算的所有精确小数。我们在交叉乘积和偏差平方的计算中，使用了货币供应量增长率的平均值 $0.560\,8/6 = 0.093\,5$ ，近似到小数点后四位。并且类似地，我们在这些计算中使用了平均通货膨胀率近似为 $0.060\,3$ 。如表中所示，我们将方差平方根的标准差的计算保留到小数点后六位。如果我们在所有的计算中使用所有精确小数，那么表中某些格子中的数值将会有所不同，而且相关系数的计算值将为 0.870 9，而不是 0.870 2，但这并不对我们的结果有本质的影响。

限的且为常数，那么相关系数就可以被有效地计算出来。后面，我们将表明当这些假设不正确时，两个不同变量间的相关系数将极大地依赖于所使用的样本。

8.2.4 相关性分析的局限

相关系数度量了两个变量间的线性关系，但它可能不总是可信的。两个变量可能具有较强的非线性关系（nonlinear relation），而且同时具有一个很低的相关系数。例如，关系 $B = (A - 4)^2$ ，不同于线性关系 $B = 2A - 4$ ，是一个非线性关系。变量 A 和 B 之间的非线性关系表现在图 8-5 中。在 A 等于 4 水平之下，变量 B 随着 A 的值的增加而减少。然而，当 A 等于 4 或者更大时，当 A 增加时， B 也增加。即使这两个变量完全相关，它们之间的相关系数也将为 0。[⊖]

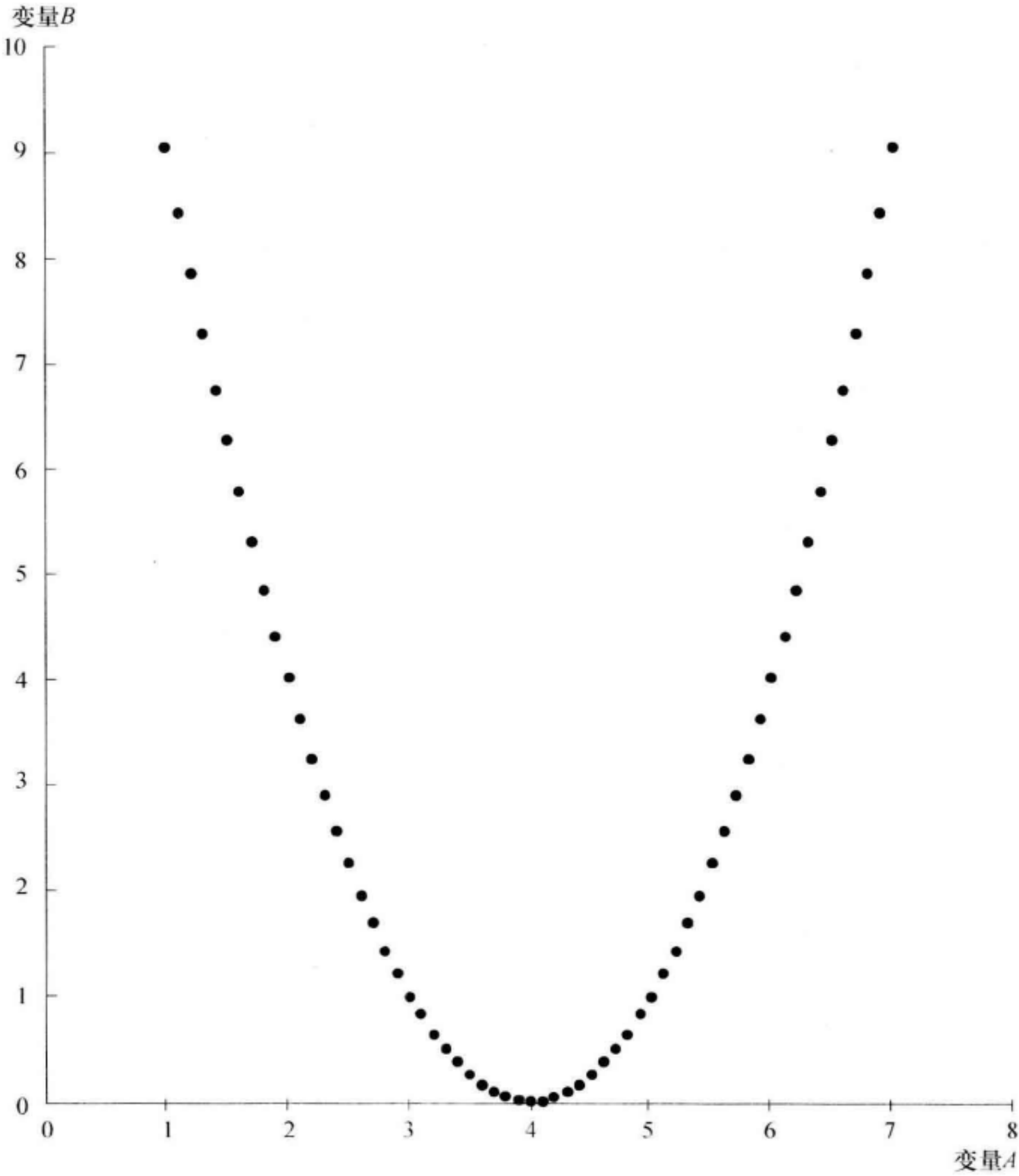


图 8-5 较强非线性关系的两个变量

当异常点出现在一个或者多个序列中时，相关系数也可能是一个不可靠的度量指标。异常点是在样本中极端的（过小或者过大）少数观测值。图 8-6 给出了在 20 世纪 90 年代（1990 年 1 月

⊖ 完全相关是指平方关系 $B = (A - 4)^2$ 。

到1999年12月)标准普尔500指数月度收益率以及美国月度通货膨胀率的散点图。

在图8-6的散点图中,大部分的数据都聚集在一起,这反映了两个变量间较不明显的关系。然而在3个数据中(3个圈出的观测值),通货膨胀率在某个特定月份大于0.8%,股票收益率极端为负。这些观测值是异常值。如果我们计算整个数据样本的相关系数,相关系数为-0.2997。如果我们消除这3个异常值,那么相关系数为-0.1347。

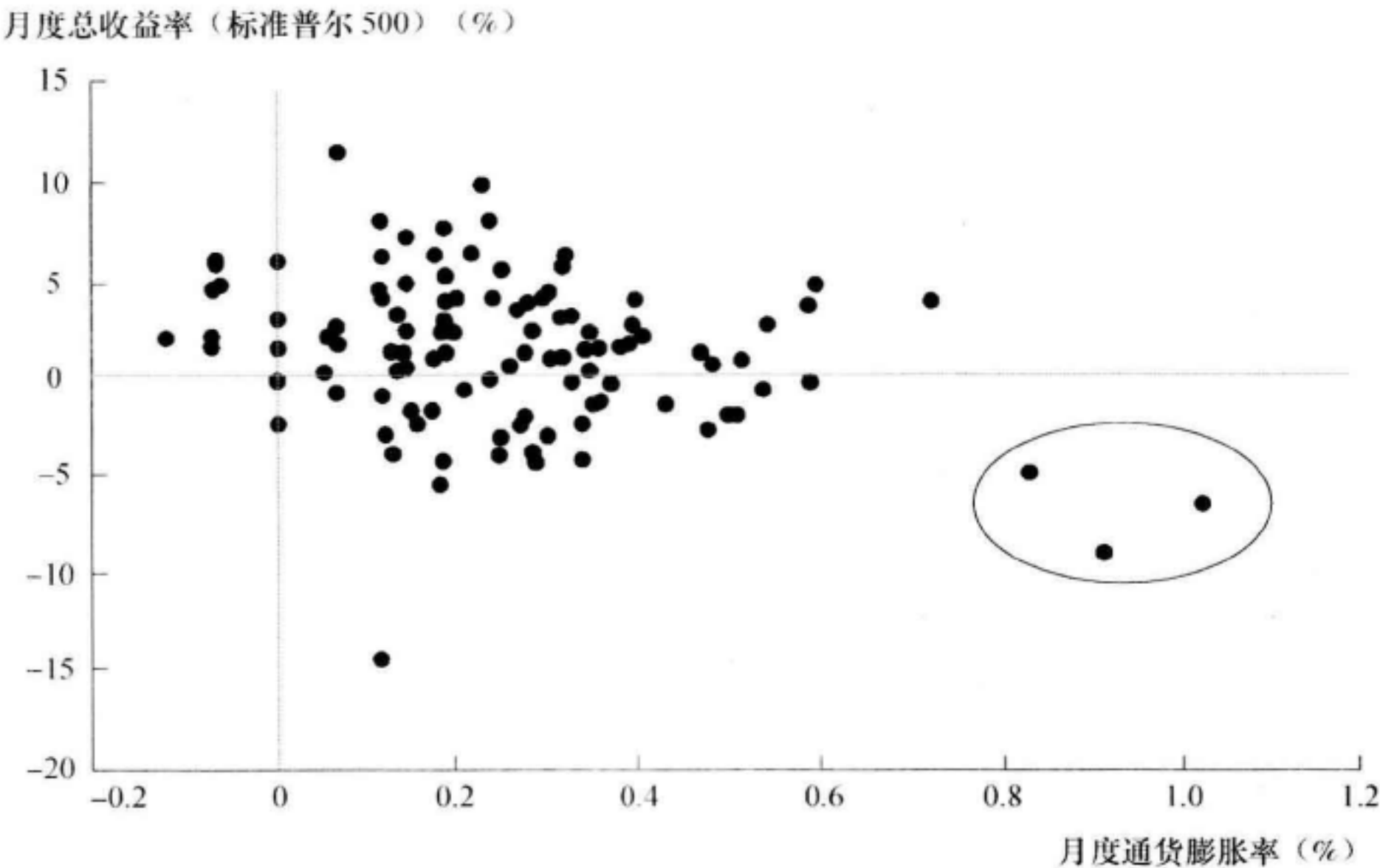


图8-6 20世纪90年代美国通货膨胀率和股票收益率

资料来源:伊博森协会(Ibboyson Associates)。

图8-6中的相关系数对于排除这3个观测值比较敏感。那么我们是否有充分的理由排除这几个观测值呢?它们是噪音还是有用的信息?一个对于图8-6部分的可能解释为在20世纪90年代,当通货膨胀率在一个月內很高时,市场参与者会关注美联储将提高利率的情况,从而造成股价下跌。这个情况提供了一个对于投资者如何对于公布高通胀消息作出反映的具有说服力的解释。结果,异常点可能提供在该时间段关于市场反映的重要信息。因此,包含异常点的相关系数可能比不包含异常点的相关系数更有意义。

作为一个一般的法则,我们必须确定是否一个计算出的样本相关系数会在移除异常点后产生巨大的变化。但是我们也必须利用我们的判断来确定那些异常点是否包含关于两个变量关系的信息(因此,异常点应该包含在相关性分析之中),或者没有包含任何信息(因此,异常点应该被排除在外)。

记住:相关关系并不意味着因果关系。即使两个变量是高度相关的,一个变量的变化也不会必然引起另一个变量的变化,也可以认为一个变量的特定值并不会造成另一个变量产生另一个特定值。进一步来说,相关关系可能在其错误地指出两个变量间关系的意义下是值得怀疑的。

术语伪相关关系(spurious correlation)是用来指:(1)两个变量间的相关关系反映某个特定数据集的偶然关系;(2)由计算所导出的相关关系将两个变量与另一个(第3个)变量混合在一起;(3)两个变量间的相关关系不是来自这两个变量的直接关系,而是由于它们与另一个(第3个)变量有关。作为第2种伪相关关系的一个例子,如果将两个不相关的变量都除以第3个变量,那么这两个变量将相关。作为第3种伪相关关系的一个例子,一个人的身高可能与一个人的词汇量大小正相关,但是这一关系是基于年龄和身高间的关系以及年龄和词汇量之间的关系所得到的。投资专家可能建议一个貌似有利可图的投资策略,实际上如果实施它的话,并不能够获利。

8.2.5 相关性分析的应用

在本节中，我们将给出一些对于实际投资问题的相关性分析。因为投资者对于通货膨胀的期望在确定资产价格时非常重要，所以通货膨胀预测的准确性将作为我们的第 1 个例子。

例 8-1 经济预测的评估 (1)

投资者密切关注经济学家对于通货膨胀的预测，但是这些预测是否包含有用的信息？在美国，职业预测家调查 (SPF) 搜集了职业预测家对于许多经济变量的预测数据。[⊖] 自从 20 世纪 80 年代早期，SPF 已经收集了美国通货膨胀率（以涵盖所有城市消费者和物品的美国消费者物价指数的变动来作为该通货膨胀率的度量指标）的预测值。如果这些关于通货膨胀的预测可以完全预测出实际的通货膨胀率，那么预测值和通货膨胀率之间的相关系数应该为 1；即预测和实际的通货膨胀率应该是相同的。

图 8-7 给出了从 1983 年第一季度到 2002 年第四季度，相对于前一个季度的当前季度 CPI 的百分比变化率的平均预测值和以年为基准的实际 CPI 百分比变化的散点图。[⊗] 在这张散点图上，x 轴反映了每个季度的预测值，而 y 轴反映了实际 CPI 的变化值。

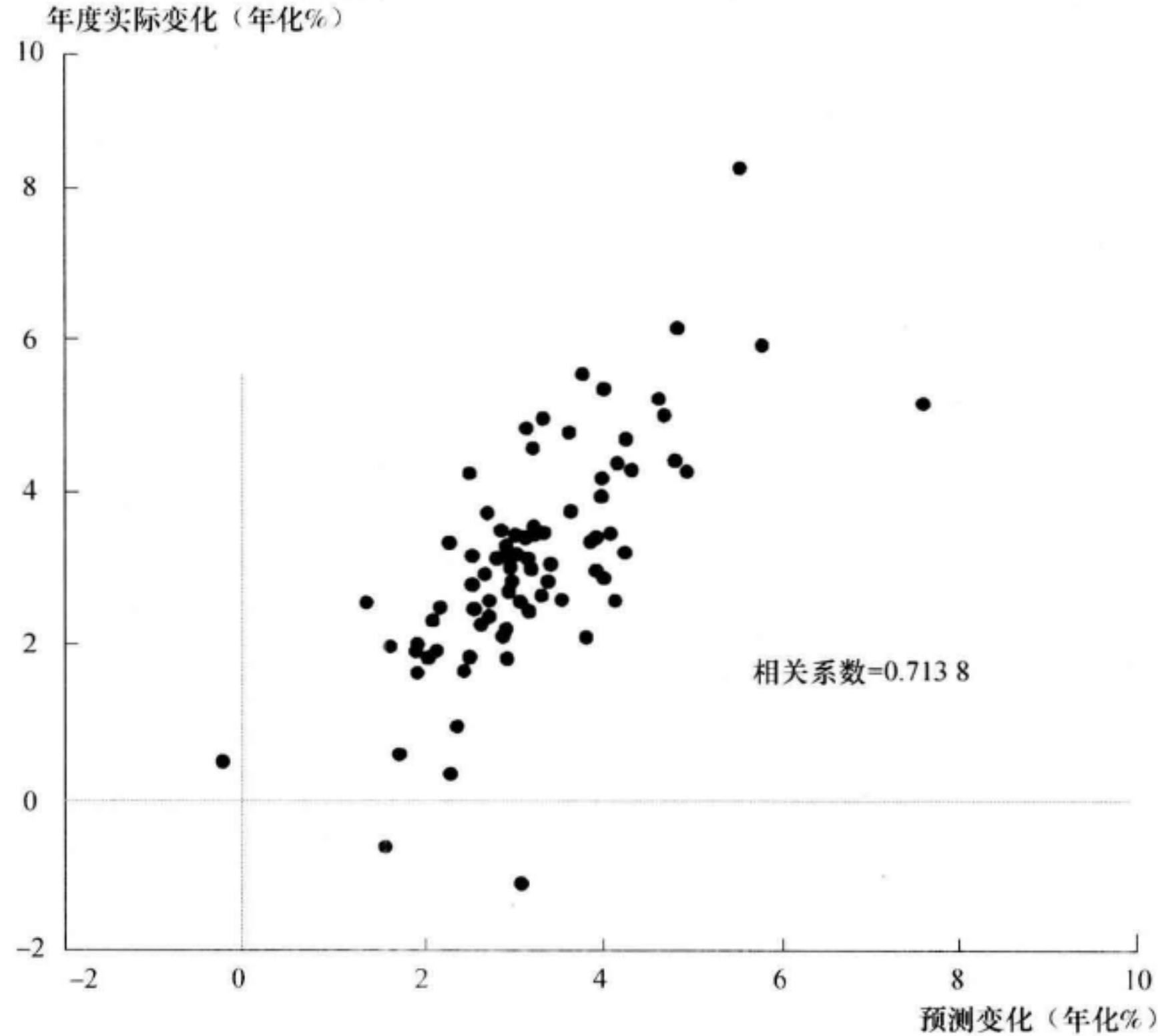


图 8-7 CPI 年度实际变化 vs 预测变化

资料来源：费城和圣路易斯联邦储备银行。

⊖ 调查最初是由 Victor Zarnowitz 为美国统计协会和国家经济研究局所进行的。从 1990 年开始，该调查由费城的联邦储备银行的 Dean Croushore 所领导执行。

⊗ 在这个散点图中，实际 CPI 的变化数据来源于美联储的经济和金融数据库，该数据可以在圣路易斯的联邦储备银行网站上获得。

如图8-7所示,在预测值和实际通货膨胀率之间存在着一个相当强的线性关系,这表明对于通货膨胀率的专业预测可能在投资决策中非常有用。事实上,这两组序列的相关系数为0.7138。虽然这里不存在因果关系,但是因为预测家们利用了信息来预测通货膨胀率,所以存在着一个直接的关系。

一个在评估投资组合经理的业绩表现中的重要问题是确定相对于该经理业绩表现的适当的基准。在最近几年,风格分析已经成了一个基准选择的重要因素。^①

例8-2 风格分析相关关系

假设一位投资组合经理在投资组合中使用了小市值的股票。通过应用风格分析,我们可以试图确定该投资组合经理是否使用了小市值增长型风格或者小市值价值型风格。

在美国,罗素2000增长型指数和罗素2000价值型指数被分别用来作为小市值增长型和小市值价值型投资经理的基准。然而,相关性分析表明这两个指数的收益率间密切相关。对截至2002年为止的20年(1983年1月到2002年12月),罗素2000增长型指数和罗素2000价值型指数的月度收益率间的相关系数为0.8526。

如果两个指数收益率间的相关系数为1,那么小市值价值型和小市值增长型股票的管理风格间将没有任何差别。如果我们知道一种风格的收益率,那么我们就一定会知道另一种风格的收益率。因为两个指数的收益率是高度相关的,所以我们可以认为两个收益率序列间存在非常小的差别。因此,我们可能无法完全区分小市值增长型和小市值价值型这两种不同的投资风格。

本章之前的例子已经讨论了两个变量间的相关关系。然而,投资经理经常需要了解许多资产收益率间的相关关系。例如,具有对于汇率变动敏感头寸的投资者必须了解不同外汇收益率和其他资产间的相关系数,以确定他们的最优投资组合和对冲策略。^②在下面的例子中,我们将会看到当我们具有两个以上的变量时,相关系数矩阵是如何来表现变量两两之间相关关系的。我们也将看到对投资经理的一个重要挑战:投资收益率的相关系数可能随着时间发生重大改变。

例8-3 汇率收益率的相关关系

汇率收益率度量的是相对于持有国外货币的国内货币区间收益率。假设英国的通货膨胀率发生了一个变动,而导致美国的美元相对于英镑的价格从1.50美元变化到了1.25美元。如果这个变动发生在某个月份,那么在那个月份相对于持有英镑的收益率将为 $(1.25 - 1.50) / 1.50 = -16.67\%$,以美元为单位。

表8-3给出了相对于持有加拿大、日本、瑞典和英国货币的美元的月度收益率相关系数矩阵。^③为了解释相关系数矩阵,我们首先讨论该表格中最上面的一个板块。这一板块第1列的数据给出的是相对于持有加元的美元收益率和相对于持有加拿大、日本、瑞典和英国货币的美元收益率间的相关系数。当然,任一个变量完全与其自身相关,因此,相对于持有加元的美元收

① 例如,参见 Sharpe (1992) 和 Buetow, Johnson 和 Runkle (2000)。

② 例如,参见 Clarke 和 Kritzman (1996)。

③ 对于20世纪80年代的数据是从1980年1月到1989年12月。而对于20世纪90年代的数据是从1990年1月到1999年12月。

益率和其自身的相关系数为 1。这一列的第 2 行给出的是 1980 ~ 1989 年间相对于持有加元的美元收益率和相对于持有日元的美元收益率之间的相关系数为 0.259 3。该板块剩余的相关系数展现的是在这个区间相对于持有其他货币组合的美元收益率的相关程度的情况。

注意表 8-3 中省略了许多相关系数。例如，该板块的第 2 列省略了相对于持有日元的美元收益率和相对于持有加元的美元收益率间的相关系数。这一相关系数被省略是因为该值与在第 1 列第 2 行的相对于持有加元的美元收益率和相对于持有日元的美元收益率间的相关系数是相同的。其他的相关系数也是由于同样的原因而被省略了。事实上，相关系数矩阵总是对称的： X 和 Y 的相关系数总是与 Y 和 X 的相关系数是相等的。

表 8-3 相对于所选的国外货币月度收益率的美元月度收益率的相关系数矩阵

1980 ~ 1989 年	加拿大	日本	瑞典	英国
加拿大	1.000 0			
日本	0.259 3	1.000 0		
瑞典	0.283 4	0.657 6	1.000 0	
英国	0.392 5	0.606 8	0.684 0	1.000 0
1990 ~ 1999 年	加拿大	日本	瑞典	英国
加拿大	1.000 0			
日本	-0.073 4	1.000 0		
瑞典	0.164 0	0.286 0	1.000 0	
英国	0.047 5	0.290 6	0.644 4	1.000 0

资料来源：伊博森协会。

如果你比较这个表格中上下两个板块，那么你将会发现许多货币间的相关系数在 20 世纪 80 年代和 20 世纪 90 年代间发生了戏剧性的变化。例如，在 20 世纪 80 年代，相对于持有日元的收益率和相对于持有瑞典克朗的收益率间的相关系数 (0.657 6)，或者相对于持有英国英镑的收益率间的相关系数 (0.606 8) 都几乎和相对于持有瑞典克朗的收益率和相对于持有英国英镑的收益率间的相关系数 (0.684 0) 一样大。然而，在 20 世纪 90 年代，日元收益率和克朗或者英镑的收益率间的相关系数都下降了超过一半 (分别降到 0.286 0 和 0.290 6)，但是克朗和英镑收益率间的相关系数却几乎没有发生变化 (0.644 4)。相对于加元的收益率和其他货币收益率间的一些相关系数甚至下降得更加厉害。在 20 世纪 80 年代，加元收益率和日元收益率间的相关系数为 0.259 3。而到了 20 世纪 90 年代，这一相关系数变为了负的 (-0.073 4)。加元的收益率和英国英镑的收益率间的相关系数从 80 年代的 0.392 5 下降到了 90 年代的 0.047 5。

最优资产配置取决于对于未来相关系数矩阵的预期。在两个资产收益率间不完全相关的正相关关系的条件下，同时持有这两种资产将获得潜在的风险降低的好处。对于未来相关系数的预期可能会基于历史样本的相关系数，但是历史相关系数的变动会给我们造成挑战。我们将在投资组合的概念一章中对这些问题进行进一步的讨论。

在下面的例子中，我们将例 8-2 中开始的对于股票市场指数相关系数的讨论延伸到对于代表大市值、小市值和广泛市场收益率的指数间相关系数的讨论。因为这些资产间的相关系数的大小告诉我们如何将资产成功地加以组合以充分分散风险，所以这一类型的分析具有非常重要的分散化和资产配置的意义。

例 8-4 不同股票收益率序列的相关系数

表 8-4 给出了 1971 年 1 月到 1999 年 12 月以及 3 个子区间（20 世纪 70 年代、80 年代和 90 年代）的 3 种美国股票指数月度收益率的相关系数矩阵。^①大市值风格是以标准普尔 500 指数的收益率来代表的，小市值风格是以空间基金咨询管理公司（Dimensional Fund Advisors）的美国小市值股票指数来代表的，而广泛市场收益率是用威尔希尔 5000 指数来代表的。

表 8-4 不同美国股票指数月度收益率的相关系数

1971 ~ 1999 年	标准普尔 500	美国小市值股票	威尔希尔 5000
标准普尔 500	1.000 0		
美国小市值股票	0.761 5	1.000 0	
威尔希尔 5000	0.989 4	0.829 8	1.000 0
1971 ~ 1979 年	标准普尔 500	美国小市值股票	威尔希尔 5000
标准普尔 500	1.000 0		
美国小市值股票	0.775 3	1.000 0	
威尔希尔 5000	0.990 6	0.837 5	1.000 0
1980 ~ 1989 年	标准普尔 500	美国小市值股票	威尔希尔 5000
标准普尔 500	1.000 0		
美国小市值股票	0.844 0	1.000 0	
威尔希尔 5000	0.991 4	0.895 1	1.000 0
1990 ~ 1999 年	标准普尔 500	美国小市值股票	威尔希尔 5000
标准普尔 500	1.000 0		
美国小市值股票	0.684 3	1.000 0	
威尔希尔 5000	0.985 8	0.776 8	1.000 0

资料来源：伊博森协会。

表 8-4 最上面板块第 1 列的数字给出几乎完全正相关的标准普尔 500 和威尔希尔 5000 收益率的相关系数：两个收益率序列间的相关系数为 0.989 4。这个结果不应该令人惊讶，因为标准普尔 500 和威尔希尔 5000 都是市值加权的指数，而且大市值的股票收益率在两个指数中的权重都很大。事实上，组成标准普尔 500 指数中的公司包含了威尔希尔 5000 指数中所有公司总市值的 80%。

小市值股票同样也与大市值股票具有相当高的相关系数。在整个样本中，标准普尔 500 收益率和美国小市值股票收益率间的相关系数为 0.761 5。而美国小市值股票收益率和威尔希尔 5000 收益率间的相关系数稍高一些（0.829 8）。这一结果同样也不令人惊讶，因为威尔希尔 5000 包含了小市值的股票，而标准普尔 500 中却没有。表 8-4 的第二、第三和第四板块给出了每个 10 年中不同股票市场收益率序列间的相关系数矩阵。例如，标准普尔 500 收益率和美国小市值股票收益率间的相关系数从 20 世纪 80 年代的 0.844 0，下降到了 90 年代的 0.684 3。^②

出于资产配置的目的，我们对于不同资产类之间相关系数的研究是为了基于预测的相关系数来保持资产适当的分散化。

① 20 世纪 70 年代数据的初始日期为 1971 年 1 月，因为威尔希尔 5000 总收益率序列是从该月份开始的。
② 20 世纪 90 年代的相关系数在 0.01 显著性水平下要小于 80 年代的相关系数。对于这类关于相关系数的假设检验可以利用费雪 z 变换的方法。关于这一方法的具体内容参见 Daniel 和 Terrell（1995）。

例 8-5 债券和股票收益率间的相关系数

表 8-5 给出了 1926 年 1 月至 2002 年 12 月间不同美国债券的收益率和标准普尔 500 的收益率的相关系数矩阵。

表 8-5 美国股票和债券收益率间的相关系数 (1926 ~ 2002 年)

所有指数	标准普尔 500	美国长期公司债券	美国长期政府债券	美国 30 天国库券	高收益公司债券
标准普尔 500	1.000 0				
美国长期公司债券	0.214 3	1.000 0			
美国长期政府债券	0.146 6	0.848 0	1.000 0		
美国 30 天国库券	-0.017 4	0.097 0	0.111 9	1.000 0	
高收益公司债券	0.647 1	0.427 4	0.313 1	0.017 4	1.000 0

资料来源：伊博森协会。

特别地，相关系数矩阵的第 1 列数字给出了标准普尔 500 收益率和不同债券收益率间的相关系数。注意这个区间中的标准普尔 500 收益率几乎完全与 30 天国库券收益率不相关 (-0.017 4)。长期公司债券收益率与标准普尔 500 收益率有一些相关关系 (0.214 3)。而高收益公司债券的收益率和标准普尔 500 收益率具有非常高的相关关系 (0.647 1)。这一高相关系数是可以理解的；因为高收益债券具有高的违约风险，所以它们的部分表现类似于股票。如果一家公司违约的话，持有高收益债券的投资者一般会损失他们所有的投资。

然而，长期政府债券和标准普尔 500 收益率间具有较低的相关系数 (0.146 6)。我们预计出这些变量间的一些相关系数，因为利率上升将会降低债券和股票未来现金流的现值。然而，两个收益率序列间相对低的相关系数表明，除了利率之外还有其他因素会影响股票的收益率。除了这些因素，债券和股票的收益率应该是高度相关的。

表 8-5 中第 2 列的数据显示，长期政府债券和公司债券之间的相关系数在这个区间中较高 (0.848 0)。虽然这一相关系数是整个矩阵中（除对角线的 1 之外）最高的，但是它不等于 1。这一相关系数小于 1 的原因是长期公司债券违约溢价会变化，而美国政府债券却不包含这一违约溢价。结果，对美国政府债券所要求收益率的变动与公司债券所要求收益率的变动间具有小于 1 的相关系数。同样注意到第 3 列中所给出的数字，高收益公司债券收益率和长期政府债券收益率间的相关系数 (0.313 1)，比高收益公司债券收益率和标准普尔 500 收益率间的相关系数要少一半。这一相对较低的相关系数从另一个角度表明，高收益债券的收益率的表现相对于债券收益率，更像股票收益率。

最后注意 30 天国库券收益率和其他所有收益率序列间具有最低的相关系数。事实上，国库券收益率和其他收益率序列间的相关系数要比在该表格中任何其他的相关系数都要小。

在本节最后的例子中，相关系数被用于财务报告的情形之中，该例子表明净利润是现金流的一个不充分的代理变量。

例 8-6 一个企业的净利润、经营现金流和自由现金流之间的相关系数

净利润 (NI)、经营现金流 (CFO) 以及自由现金流 (FCFF) 是本书作者常用的对于公司业绩表现进行评估的度量指标。如果这些指标是高度相关的话，那么对于给定的公司，这些指标间的不同将不会造成对于公司相对估值结果的差别。

CFO 等于净利润加上非现金的净支出（该支出在计算净利润时被减去），然后再减去在相同时间段中公司在运营资本中的投资额。FCFF 等于 CFO 加上税后利息支出，再减去公司在该时间段中对于固定资本的投资额。FCFF 可能被解释为在所有经营费用被支付以及运营和固定资本的必要投资做出之后的公司供应商（债权人和股东）可获得的现金流资本。^①

一些分析师只基于 NI 做出他们的估值，而忽略了 CFO 和 FCFF。如果 NI、CFO 和 FCFF 之间的相关系数很高，那么这些分析师忽略 CFO 和 FCFF 的决定是容易为我们所理解的，因为 NI 也能够包含人们所需要了解的关于现金流的所有信息。

表 8-6 给出了 2001 年 6 家美国经营女性服装上市公司的 NI、CFO 和 FCFF 间的相关系数。在计算相关系数之前，我们先将所有的数据标准化，即把每家公司的 3 个业绩表现的度量指标除以公司在该年的收入。^②

表 8-6 2001 年美国女性服装商店的业绩表现度量指标

	NI	CFO	FCFF
NI	1.000 0		
CFO	0.695 9	1.000 0	
FCFF	0.404 5	0.821 7	1.000 0

资料来源：Compustat。

因为 CFO 和 FCFF 包括 NI 作为一个成分（这指的是 CFO 和 FCFF 可以通过在 NI 的基础上加上或者减去不同的量来获得），我们预计 NI 和 CFO 以及 NI 和 FCFF 的相关系数可能为正。表 8-6 支持这一结论。然而，这些与 NI 相关的相关系数比起 CFO 和 FCFF 间的相关系数要小得多（0.821 7）。表中最小的相关系数是 NI 和 FCFF 间的相关系数（0.404 5）。这个相对较小的相关系数表明 NI 包含一些不同于 2001 年对于这些公司的 FCFF 的信息。在本章之后的部分，我们将检验 NI 和 FCFF 间的相关系数是否是显著不同于 0 的。

8.2.6 相关系数显著性检验

显著性检验使得我们能够评估随机变量之间的明显关系是否只是源于运气的结果。如果我们能确定这个关系不是来源于运气，那么我们将能够在预测中使用该信息，因为对于一个变量的好预测将帮助我们预测另一个变量。使用表 8-2 中的数据，我们计算出 0.870 2，其为 1970 ~ 2001 年 6 个工业化国家的长期货币增长和长期通货膨胀间的相关系数。这个估计出的相关系数看上去很高，但它是否显著不等于 0？在我们能够回答这个问题之前，我们必须了解一些对应变量自身分布的具体情况。为了简化问题，我们假设这两个变量服从正态分布。^③

我们提出两个假设：原假设 H_0 ，总体的相关系数为 0 ($\rho = 0$)；而备择假设 H_a ，总体的相关系数不等于 0 ($\rho \neq 0$)。

备择假设是一个关于相关系数不等于 0 的检验；因此，一个双侧的检验是恰当的。^④只要两个变量服从正态分布，我们可以利用样本相关系数 r ，通过检验来确定原假设是否应该被拒绝。该 t 检验的公式为

① 更多关于这 3 个度量指标以及它们在股票估值中的应用，参见 Stowe, Robinson, Pinto 和 McLeavey (2002)。在脚注段落中的叙述根据美国一般公认会计原则解释了这些度量指标间的关系。Stowe 等人也根据国际会计准则讨论了这些度量指标间的关系。

② 该表中的结果基于所有在 2001 年年末市值超过 250 000 000 美元的女性服装商店（美国职业安全与卫生管理局的标准工业分类 5621）的数据。市值标准是用来排除微小市值公司的，因为它们的业绩表现度量指标的相关系数可能与那些高市值的公司有明显不同。我们将在第 9 章回归的错误设定一节中具体讨论该例子中所使用的数据标准化方法（除以公司收入）。

③ 事实上，我们必须假设变量来源于一个二元正态分布。如果两个变量 X 和 Y 来自一个二元正态分布，那么对于 X 的每个值， Y 的分布为正态分布。例如，参见 Ross (1997) 或者 Greene (2003)。

④ 参见假设检验一章，以了解关于双侧检验更加深入的讨论。

$$t = \frac{r \sqrt{n-2}}{\sqrt{1-r^2}} \quad (8-3)$$

如果原假设为真的话,该检验统计量是一个自由度为 $n-2$ 的 t 分布。一个对于式 (8-3) 的实际观测为拒绝原假设 $H_0: \rho=0$ 所需要 r 的大小随着样本量 n 的增加而减少。这有两个原因。首先,随着 n 的增加,自由度的大小也会增加,而临界值 t_c 的绝对值会减少。其次,具有较大 n 值的分子的绝对值会增大,从而导致一个更大的 t 值。例如,对于样本量 $n=12$ 的样本, $r=0.58$ 将导致 t 统计量的值为 2.252,其仅在 0.05 显著性水平上 ($t_c=2.228$) 显著。对于样本量 $n=32$ 的样本,一个更小的样本相关系数 $r=0.35$ 将导致一个 t 统计量为 2.046,其仅在 0.05 显著性水平上显著 ($t_c=2.042$); $r=0.35$ 在样本量 $n=12$ 的样本中甚至在 0.10 显著性水平下也是不显著的。另一种说明这一点的方法为从一个相同的总体中进行抽样,当所有其他条件不变的情况下,一个不真的原假设 $H_0: \rho=0$ 在我们增加样本量时,更可能被拒绝。

例 8-7 货币供应量增长率和通货膨胀率间相关系数的检验

在本章的前面,我们给出了从 1970~2001 年间的 6 个工业化国家长期货币供应量增长率和通货膨胀率间的样本相关系数为 0.870 2。假设我们想要检验原假设 H_0 : 真正总体的相关系数为 0 ($\rho=0$) 相对于备择假设 H_a : 总体相关系数不为 0 ($\rho \neq 0$)。

回忆一下,该样本具有 6 个观测值,我们可以计算检验该原假设的统计量如下:

$$t = \frac{0.8702 \sqrt{6-2}}{\sqrt{1-0.8702^2}} = 3.532$$

检验统计量的值为 3.532。因为在双侧检验 t 分布的临界值数值表上给出了对于自由度为 $n-2=6-2=4$ 的 t 分布在 0.05 显著性水平下的临界值,所以如果检验统计量的值大于 2.776 或者小于 -2.776,我们就能拒绝原假设 (总体相关系数等于 0)。我们仅基于 6 个观测值而拒绝不相关原假设的情况是不经常发生的;它更加表现出这 6 个国家的长期货币供应量增长率和长期通货膨胀率间较强的相关关系。

例 8-8 瑞典克朗-日元收益率间相关系数的检验

表 8-3 中的数据给出了 1990 年 1 月至 1999 年 12 月间相对于瑞典克朗和相对于日元的美元月度收益率间的样本相关系数为 0.286 0。如果我们观测该样本相关系数,那么我们是否能够拒绝所关心的或者总体的相关系数为 0 的原假设?

对于从 1990 年 1 月到 1999 年 12 月的 120 个月份,我们使用下面的统计量来检验原假设 H_0 : 总体真实的相关系数为 0,对应于备择假设 H_a : 总体的相关系数不等于 0:

$$t = \frac{0.2860 \sqrt{120-2}}{\sqrt{1-0.2860^2}} = 3.242$$

在 0.05 显著性水平下,该检验统计量的临界值为 1.98 ($n=120$, 自由度=118)。当检验统计量或者大于 1.98 或者小于 -1.98 时,我们就会拒绝总体相关系数为 0 的原假设。检验统计量为 3.242,故我们应该拒绝原假设。

注意到即使该相关系数比之前例子中的相关系数要小,但是在该情况中的样本相关系数在 0.05 显著性水平下也显著不为 0。虽然相关系数更小,但是因为该样本量更大 (120 个观测值而不是 6 个),所以其仍旧非常显著。

上面的例子表明样本大小在相关系数显著性检验中非常重要。下面的例子也表明了样本大小的重要性,并检验了在0.01和0.05显著性水平下的相关关系。

例 8-9 债券收益率和国库券收益率间的相关系数

表8-5给出了1926年1月到2002年12月间美国政府债券月度收益率和30天国库券月度收益率间的样本相关系数为0.1119。假如我们想要检验该相关系数是否显著不为0。从1926年1月到2002年12月共有924个月。因此,为了检验原假设 H_0 (总体真实的相关系数为0)对应于备择假设 H_a (总体的相关系数不为0),我们使用下面的检验统计量:

$$t = \frac{0.1119 \sqrt{924-2}}{\sqrt{1-0.1119^2}} = 3.4193$$

在0.05显著性水平下,该检验量的临界值近似为1.96。在0.01显著性水平下,该检验量的临界值近似为2.58。该检验量为3.4193,因此我们在0.05和0.01显著性水平下都会拒绝总体间没有相关性的原假设。该例子表明在大样本的情况下,即使相对较小的相关系数也会表现出显著不为0的结果。

在本节最后一个例子中,我们将讨论小样本的另一种情况。

例 8-10 公司净利润和自由现金流之间相关系数的检验

在本章的前面,我们曾给出2001年6家女性服装商店的NI和FCFF之间的样本相关系数为0.4045。假如我们想要检验原假设 H_0 :总体真实的相关系数为0($\rho=0$)对应于备择假设 H_a :总体的相关系数不为0($\rho \neq 0$)。回忆一下,该样本具有6个观测值,我们可以计算检验该原假设的统计量如下:

$$t = \frac{0.4045 \sqrt{6-2}}{\sqrt{1-0.4045^2}} = 0.8846$$

在自由度为 $n-2=6-2=4$ 以及0.05的显著性水平下,如果检验统计量的值大于2.776或者小于-2.776,我们会拒绝总体相关系数为0的原假设。然而,在这个例子中,该 t 统计量为0.8846,所以我们无法拒绝原假设。因此,对于女性服装商店的这个样本,当每个指标除以公司销售量进行标准化之后,NI和FCFF之间不存在统计上显著的相关性。^①

散点图给出了两个变量间关系的可视化图像表示,而相关系数给出了任一线性关系的数量化表示。相关系数的绝对值越大,说明线性关系越强。正的相关系数表明了两个数据集的一个正向的相关变动关系,而负的相关系数表明了一个相反的变动关系。在例8-8和8-9中,我们看到相对较小的样本相关系数(0.2860和0.1119)也可能表现出统计上的显著,因此,它们也能提供关于经济变量变动行为的有价值信息。

下面我们将介绍线性回归——在检验两个变量间关系中很重要的另一个工具。

① 这里值得再重复一下,当样本越小,以拒绝相关系数为0的原假设所需的样本相关系数数值为证据的值就要越大。在样本量为6的情况下,样本相关系数的绝对值需要大于0.81(保留到小数点后两位)才能使我们拒绝原假设。从另一方面来看,如果样本量为24,那么文中的0.4045这个值将会变得显著,因为 $0.4045(24-3)^{1/2}/(1-0.4045^2)^{1/2}=2.075$,其比0.05显著性水平下自由度为22的临界 t 值2.074要大。

8.3 线性回归

8.3.1 单变量的线性回归

作为一名金融分析师，你会经常想要了解金融或者经济变量间的关系，或者想要利用关于一个变量的信息来预测另一个变量的值。例如，你可能想要知道 10 年期国库券收益率的变动对于标准普尔 500 净收益率的影响（净收益率是市场价格与公司收益比值的倒数）。如果这两个变量间的关系是线性的，那么你就可以使用线性回归来描述它。

线性回归使我们能用一个变量来预测另一个变量、检验两个变量间的关系以及对于两个变量间相关程度给予定量表示。本章的余下部分将关注于单个自变量的线性回归。在下一章中，我们将考察多于一个自变量的回归。

回归分析是从你所要试图解释的因变量（记为 Y ）开始讨论的。而自变量（记为 X ）是你用来解释因变量变化的变量。例如，你可能会基于标准普尔 500 指数（自变量）来试图解释小盘股的收益率（因变量）。或者你可能试图用一个国家货币供应量增长率（自变量）的函数来试图解释通货膨胀率（因变量）。

线性回归（linear regression）假设因变量和自变量之间存在着某种线性关系。下面的回归方程描述了这一关系：

$$Y_i = b_0 + b_1 X_i + \varepsilon_i, i = 1, \dots, n \quad (8-4)$$

这个方程表明因变量（dependent variable） Y 等于截距 b_0 加上斜率系数 b_1 乘以自变量（independent variable） X ，再加上误差项（error term） ε 。误差项代表因变量中无法被自变量所解释的部分。我们将截距 b_0 和斜率 b_1 称为回归系数（regression coefficients）。

回归分析使用两种主要的数据类型：横截面数据和时间序列数据。横截面数据涉及相同时间区间上许多 X 和 Y 的观测值。这些观测值可能来自不同的公司、资产类、投资基金、人群、国家或者其他与回归模型相关的实体。例如，一个横截面模型可能会使用许多公司的数据来检验在特定时间区间中是否预测的每股盈利增长率能够解释市盈率间的差别。“解释”这个词在描述回归关系时会经常被使用。公司 P/E 的一个不依赖于其他任何变量的估计值为平均的 P/E。如果关于一个自变量对于 P/E 的回归倾向于给出比假设公司的 P/E 等于平均的 P/E，更加精确的 P/E 估计值的话，我们就认为自变量能够帮助我们预测 P/E，因为利用自变量改善了我们的预测。最后，注意到如果我们在回归中使用横截面数据，我们常将这些观测值记为 $i = 1, 2, \dots, n$ 。

时间序列数据使用的是不同时间点上对于相同公司、资产类、投资基金、个人、国家或者其他与回归模型相关的实体的观测值。例如，一个时间序列模型可能使用许多年的月度数据来检验美国的通货膨胀率是否决定了美国的短期利率。^① 如果我们在回归中使用时间序列数据，我们经常将观测值记为 $t = 1, 2, \dots, T$ 。^②

那么应该如何恰当地通过线性回归模型来预测 b_0 和 b_1 呢？线性回归也被称为线性最小二乘法，它会计算出与观测值拟合最好的那条直线；它选择截距 b_0 和斜率 b_1 的值以使得观测值和回

① 时间序列和横截面的混合数据，也被称为面板数据，当前在金融分析中常常被使用。面板数据的分析是一个高级的话题，Greene（2003）对于该内容进行了详细的讨论。

② 在本章中，我们主要使用记号 $i = 1, 2, \dots, n$ ，甚至对于时间序列也是这样，这样做是为了避免前后不同记号的切换所可能造成的麻烦。

归直线间纵向距离的平方和最小。线性回归选择式 (8-4) 中的估计值 (estimated) 或者拟合值 (fitted parameters) $\hat{b}_0$ 和 $\hat{b}_1$ 以使得下式达到最小:[Ⓐ]

$$\sum_{i=1}^n (Y_i - \hat{b}_0 - \hat{b}_1 X_i)^2 \tag{8-5}$$

在这个式子中, $(Y_i - \hat{b}_0 - \hat{b}_1 X_i)^2$ 这项表示 (因变量 - 因变量的预测值)²。通过利用这一方法估计出 $\hat{b}_0$ 和 $\hat{b}_1$ 的值之后, 我们就可以拟合出一条通过 X 和 Y 观测值的直线, 该直线能够在给定任一 X 值的情况下, 最好地解释 Y 的值。[Ⓑ]

注意到我们无法在回归模型中观测到总体参数值 b_0 和 b_1 。相反, 我们只能观测到 $\hat{b}_0$ 和 $\hat{b}_1$, 它们是总体参数的估计值。因此, 预测必须要基于参数的估计值, 而且关于假定总体参数值的检验也要基于这些参数的估计值来进行。

图 8-8 给出了线性回归是如何表现的一个可视化的例子。该图向我们展现了由对于 1970 ~ 2001 年 6 个工业化国家 ($n=6$) 的年度通货膨胀率 (因变量) 和年度货币供应量增长率 (自变量) 的回归关系进行估计而得到的线性回归的结果。[Ⓒ]所估计的方程为:

长期通货膨胀率 = $b_0 + b_1$ (长期货币供应量增长率) + ε

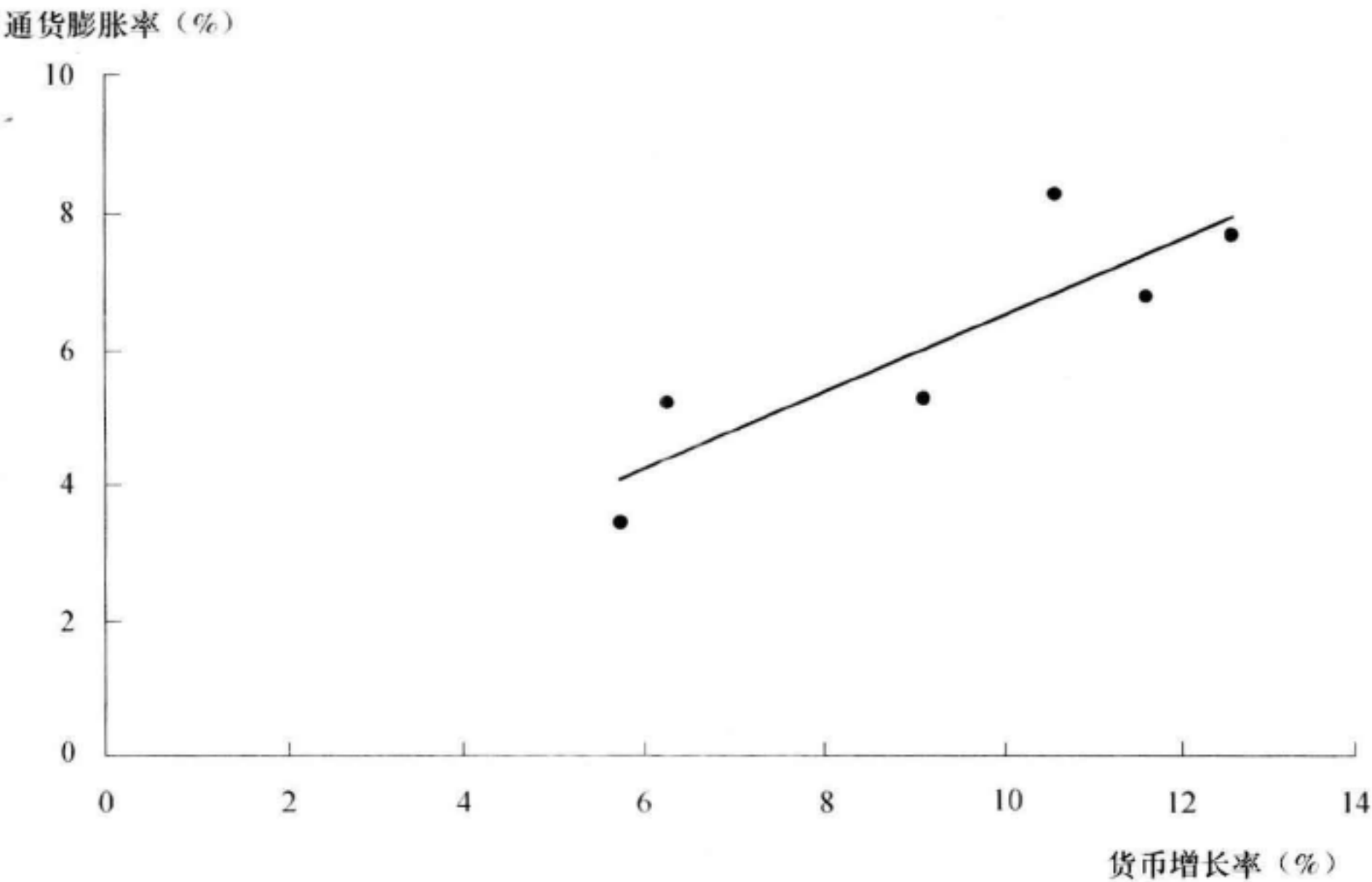


图 8-8 用各国货币供应量的增长率来解释通货膨胀率的拟合回归直线 (1970 ~ 2001 年)
资料来源: 国际货币基金。

6 个数据点中的每一个与拟合回归直线间的距离是回归的残差, 其等于因变量的实际值与由回归方程给出的因变量的预测值之间的差值。线性回归会选择式 (8-4) 中使纵向距离的平方和最小的 $\hat{b}_0$ 和 $\hat{b}_1$ 作为估计系数。所估计出的回归方程为长期通货膨胀率 = 0.008 4 + 0.554 5 (长期货币供应量增长率)。[Ⓓ]

根据该回归方程, 如果对于任一国家的长期货币供应量增长率为 0, 那么这个国家的长期通

Ⓐ 系数符号上的“帽子”表明该值是一个估计量。
Ⓑ 精确的关于最优估计量 b_0 和 b_1 的统计学概念的讨论, 参见 Greene (2003)。
Ⓒ 表 8-2 给出了这些数据。
Ⓓ 我们将月度收益率以小数的形式代入方程。

货膨胀率将为 0.84%。对于任一国家长期货币供应量增长率的每一个百分点的增长,长期通货膨胀率将预期增长 0.554 5 个百分点。在如同这个回归,即包含一个自变量的回归中,其斜率等于 $\text{Cov}(Y, X)/\text{Var}(X)$ 。我们能利用表 8-2 中的数据计算出其回归方程的斜率,具体过程摘录如下:

表 8-2 (摘录)

	货币供应量增长率 X_i	通货膨胀率 Y_i	交叉乘积 $(X_i - \bar{X})(Y_i - \bar{Y})$	平方偏差 $(X_i - \bar{X})^2$	平方偏差 $(Y_i - \bar{Y})^2$
总和	0.560 8	0.361 6	0.002 185	0.003 939	0.001 598
平均值	0.093 5	0.060 3			
协方差			0.000 437		
方差				0.000 788	0.000 320
标准差				0.028 071	0.017 889

$$\begin{aligned}\text{Cov}(Y, X) &= 0.000\,437 \\ \text{Var}(X) &= 0.000\,788 \\ \text{Cov}(Y, X)/\text{Var}(X) &= 0.000\,437/0.000\,788 \\ \hat{b}_1 &= 0.5545\end{aligned}$$

在一个线性回归中,拟合的回归直线会通过因变量的均值和自变量的均值这一点。如同表 8-1 所示(摘录如下),从 1970~2001 年这 6 个国家的平均长期货币供应量增长率为 9.35%,而平均长期通货膨胀率为 6.03%。

表 8-1 (摘录)

	货币供应量增长率	通货膨胀率
平均值	9.35%	6.03%

因为点 (9.35, 6.03) 在回归直线 $\hat{b}_0 = \bar{Y} - \hat{b}_1 \bar{X}$ 上,所以我们可以利用该点来求解其截距,具体过程如下:

$$\hat{b}_0 = 0.060\,3 - 0.554\,5(0.093\,5) = 0.008\,4$$

我们上面给出了如何来求解线性回归方程的每一步过程,这样,我们能十分清晰地了解每个数据是怎么得到的。通常,一位分析师会使用电子表格或者统计软件包中的数据分析函数来进行线性回归分析。后面,我们将讨论如何利用回归残差来对于一个回归模型的不确定性进行定量分析。

8.3.2 线性回归模型的前提假设

我们已经讨论了如何解释一个线性回归模型中的系数。现在我们转向对于该模型的统计学假设。假设我们具有 n 个关于因变量 Y 和自变量 X 的观测值,而且我们想要估计式 (8-4):

$$Y_i = b_0 + b_1 X_i + \varepsilon_i, i = 1, \cdots, n$$

为了从单个自变量的线性回归模型中获得有效的结论,我们需要做出如下 6 个假设,它们被称为经典正态线性回归的假设:

- 1. 因变量 Y 和自变量 X 之间的关系是关于参数 b_0 和 b_1 线性的。这一要求意味着 b_0 和 b_1 的次数是一次的, b_0 和 b_1 都没有和另一个回归参数相乘或者相除 (例如, b_0/b_1)。这个要求并没有要求 X 的次数一定要等于 1。
- 2. 自变量 X 不是随机的。[⊖]

⊖ 虽然我们假设回归模型中的自变量不是随机的,但是这个假设通常是不满足的。例如,假设标准普尔 500 的月度收益率不是随机的是不现实的。如果自变量是随机的,那么回归模型是否会不正确呢? 幸运的是,回答是否定的。计量经济学家已经证明即使自变量是随机的,我们仍旧可以利用给定误差项和自变量不相关这个关键假设下的回归模型的结果。然而,对于该可信程度的数学证明是有相当难度的。例如,参见 Greene (2003) 或者 Goldberger (1998)。

3. 误差项的期望值为0: $E(\varepsilon) = 0$ 。
4. 误差项的方差对于所有观测值都是相同的: $E(\varepsilon_i^2) = \sigma_\varepsilon^2, i = 1, \dots, n$ 。
5. 误差项 ε 是和观测值不相关的。因此, 对于所有不相等的 i 和 j 来说, $E(\varepsilon_i \varepsilon_j) = 0$ 。^①
6. 误差项 ε 服从正态分布。^②

现在我们可以进一步深入对于每一个假设进行考察。

假设1 对于一个有效的线性回归非常重要。如果自变量和因变量之间的关系关于参数是非线性的, 那么用线性回归模型对于该关系进行估计将会得到错误的结果。例如, $Y_i = b_0 e^{b_1 X_i} + \varepsilon_i$ 就是关于 b_1 非线性的, 因此我们无法对其应用线性回归模型。^③

即使因变量是非线性的, 只要回归是关于参数线性的, 我们仍可以使用线性回归。例如, 线性回归可以用来估计等式 $Y_i = b_0 + b_1 X_i^2 + \varepsilon_i$ 。

假设2 和3 保证了线性回归会得到正确的估计值 b_0 和 b_1 。

假设4、5 和6 使我们能利用线性回归模型来确定估计参数 $\hat{b}_0$ 和 $\hat{b}_1$ 的分布, 从而使我们能检验这些系数是否等于某个特定值。

- 假设4, 误差项的方差对于所有的观测值都是相等的, 这也被称为同方差性假设。回归分析一章将会讨论对于这个假设的检验以及在该假设被违背情况下的正确处理方法。
- 假设5, 误差项和观测值不相关, 这也是对于正确获得出估计参数 $\hat{b}_0$ 和 $\hat{b}_1$ 的方差来说非常必要的条件。多元回归一章将会对于该假设不满足的情况进行讨论。
- 假设6, 误差项服从正态分布, 该假设也使得我们能够简单地检验关于一个线性回归模型的特定假设。^④

例 8-11 经济预测的评估 (2)

如果经济预测是完全准确的, 那么对于某个季度经济变量变化的每个预测都会准确地与那个季度实际的变动值相一致。即使预测是不准确的, 我们也希望至少它们是无偏的, 即预测误差的期望值等于0。一个无偏的预测可以表达为 $E(\text{实际变动值} - \text{预测变动值}) = 0$ 。事实上, 大部分对于预测准确性的评估所检验的是预测值是否是无偏的。^⑤

图8-9 重复了图8-7 的结果, 给出了从1983年第一季度到2002年第四季度, 以年化值作为单位的, 相对于前一个季度的当前季度CPI百分比变化的平均预测值和实际的CPI百分比变动值, 但是图8-9增加了对于等式: 实际的百分比变动值 = $b_0 + b_1(\text{预测的百分比变动值}) + \varepsilon$ 的拟合

① $\text{Var}(\varepsilon_i) = E[\varepsilon_i - E(\varepsilon_i)]^2 = E(\varepsilon_i - 0)^2 = E(\varepsilon_i^2)$ 。 $\text{Cov}(\varepsilon_i, \varepsilon_j) = E\{[\varepsilon_i - E(\varepsilon_i)][\varepsilon_j - E(\varepsilon_j)]\} = E[(\varepsilon_i - 0)(\varepsilon_j - 0)] = E(\varepsilon_i \varepsilon_j) = 0$ 。

② 如果回归的误差项不服从正态分布, 那么我们仍可以使用回归分析。计量经济学家利用 χ^2 检验而不是 F 检验来对假设进行检验, 从而摆脱了这个正态性假设。这一差别一般不会导致某个检验会拒绝一个特定的原假设的结果。

③ 关于参数非线性的更多内容, 参见 Gujarati (2003)。

④ 对于大样本来说, 我们也许可以因为抽样一章中所讨论的中心极限定理, 而舍弃该正态性假设; 参见 Greene (2003)。渐进理论表明在许多情况中, 由标准回归程序所得到的检验统计量即使在误差项不服从正态分布的情况下也是有效的。然而, 正如统计学概念和市场收益率一章中所阐述的, 一些金融时间序列的非正态性可能会非常严重。在非正态情况非常严重的情况下, 即使拥有相对较多的观测值, 希望藉由渐近理论的证明而使用从线性回归模型所得到的检验统计量, 也许是不恰当的。

⑤ 例如, 参见 Keane 和 Runkle (1990)。

回归直线。如果预测是无偏的，那么截距 b_0 应该等于 0，而 b_1 应该等于 1。我们也应该发现 $E(\text{实际变动值} - \text{预测变动值}) = 0$ 。如果预测值实际上是无偏的，只要 $b_0 = 0$ 和 $b_1 = 1$ ，那么误差项 $[\text{实际变动值} - b_0 - b_1(\text{预测变动值})]$ 的期望值将为 0，正如线性回归模型的假设 3 所要求的。对于无偏的预测值， b_0 和 b_1 的任何其他值都将产生一个期望值不等于 0 的误差项。

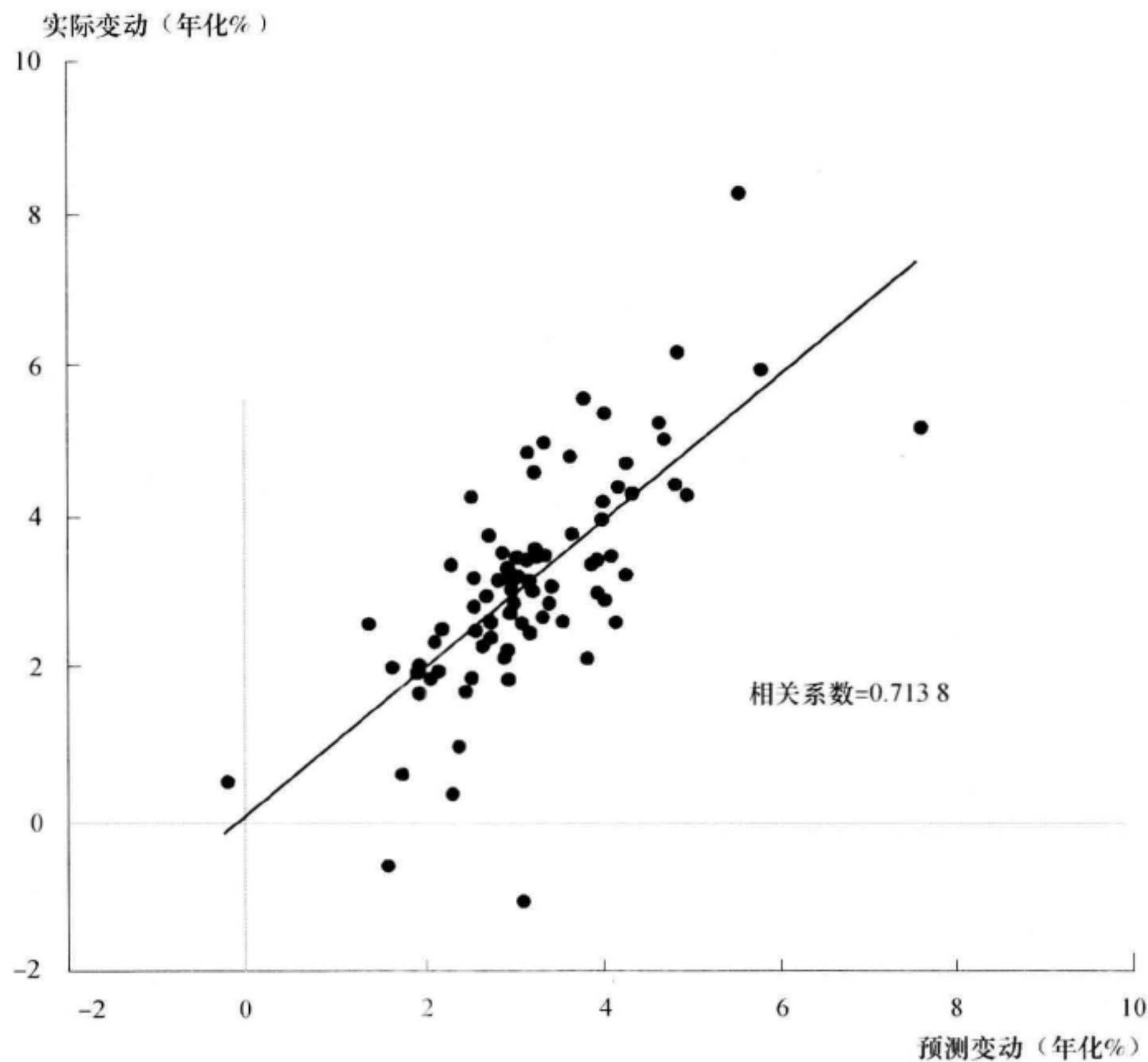


图 8-9 CPI 的实际变动相对于预测变动
资料来源：费城和圣路易斯联邦储备银行。

如果 $b_0 = 0$ 而 $b_1 = 1$ ，如果职业预测家对于 CPI 变化的预测为 0，那么我们对于 CPI 实际变化的最佳猜测为 0。对于职业预测家每一百分点变动预测的增加，回归模型将预测出一个百分点实际变动的增加。

图 8-9 中的拟合回归直线来源于等式：实际变动 = $-0.0140 + 0.9637(\text{预测变动})$ 。注意到 b_0 和 b_1 的预测值都接近于与无偏预测相一致的值 $b_0 = 0$ 和 $b_1 = 1$ 。在本章的后面，我们将讨论如何检验 $b_0 = 0$ 和 $b_1 = 1$ 的假设。

8.3.3 估计量的标准误差

线性回归模型有时能较好地描述两个变量间的关系，但是有时却不能。我们必须能够区分这

两种情况以有效地使用回归分析。因此，在这一节和之后的几节中，我们将讨论如何度量一个给定线性回归模型描述因变量和自变量之间关系的好坏。

例如，图 8-9 显示预测的通货膨胀率和实际的通货膨胀率之间具有一个较强的线性关系。如果我们了解到职业预测家对于一个特定季度通货膨胀率的预测，那么我们将有理由确信我们能相对准确地通过该回归模型来预测实际的通货膨胀率。

然而，在其他的情况中，因变量和自变量之间的关系并不强。图 8-10 在图 8-6 的 20 世纪 90 年代通货膨胀率和股票收益率的数据基础上添加了一条拟合的回归直线。在这张图上，实际的观测值比图 8-9 要更远离拟合的回归直线。假设一个特定的通货膨胀水平，利用该估计的回归方程来预测月度股票收益率会导致不准确的预测。

正如所指出的，图 8-10 的回归关系没有图 8-9 精确。估计量的标准误（有时也称为回归的标准误）度量了这一不确定性。除了它度量的是 $\hat{\varepsilon}_i$ （回归的残差项）的标准差之外，该统计量非常像单个变量的标准差。

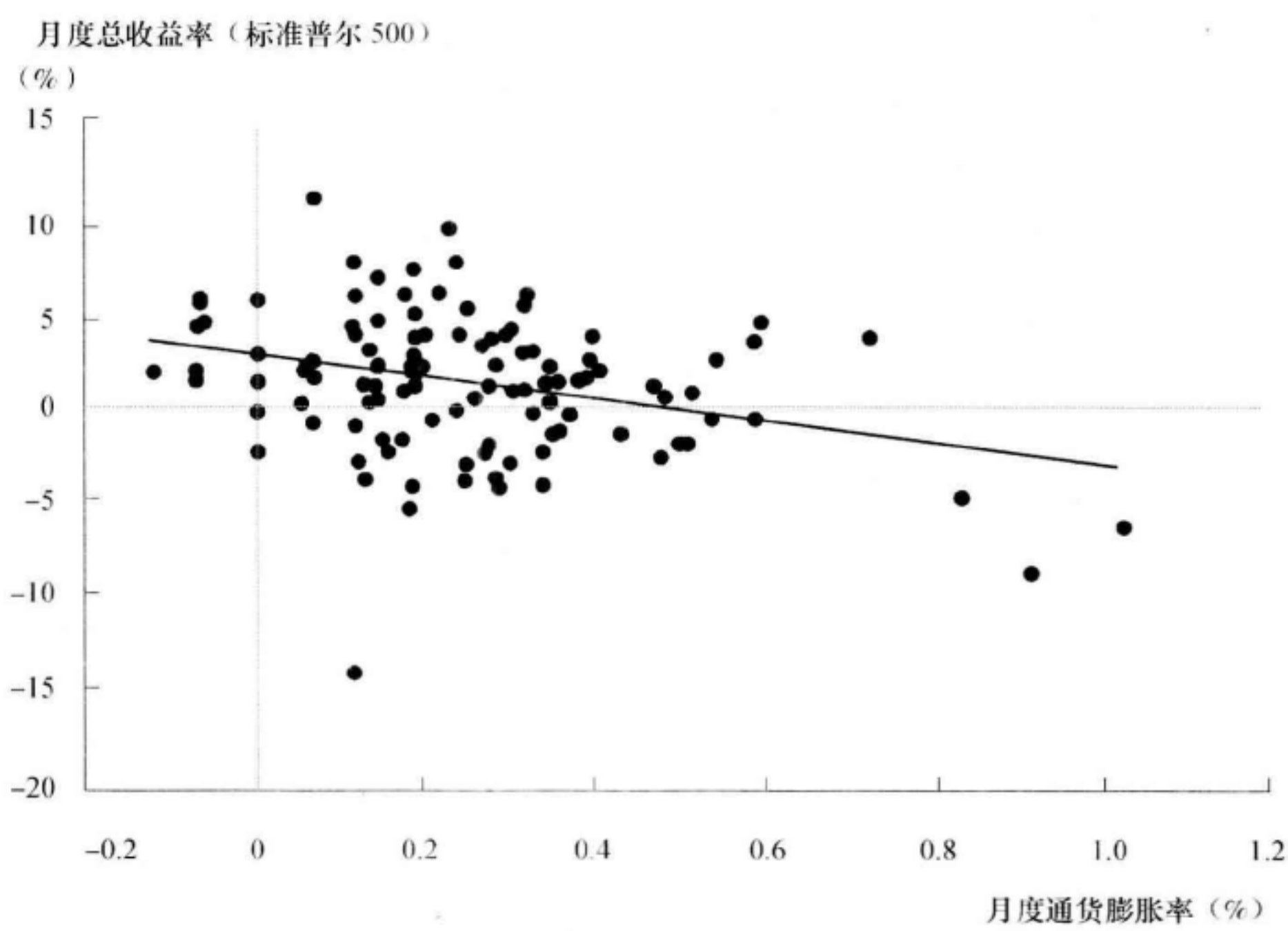


图 8-10 由通货膨胀率解释 20 世纪 90 年代股票价格的拟合回归直线
资料来源：伊博森协会。

单个自变量的线性回归模型的估计量标准差（SEE）的公式为

$$SEE = \left(\sum_{i=1}^n \frac{(Y_i - \hat{b}_0 - \hat{b}_1 X_i)^2}{n - 2} \right)^{1/2} = \left(\sum_{i=1}^n \frac{(\hat{\varepsilon}_i)^2}{n - 2} \right)^{1/2} \tag{8-6}$$

在该等式的分子中，我们计算了每个观测值因变量的实际值和其预测值 $(\hat{b}_0 + \hat{b}_1 X_i)$ 之间的差值。因变量的实际值和预测值之间的差值是回归的残差 $\hat{\varepsilon}_i$ 。

式（8-6）看上去非常像计算标准差的公式，除了分母为 $n - 2$ 而不是 $n - 1$ 。我们使用 $n - 2$ 是因为样本包含 n 个观测值，而且线性回归模型有两个估计参数（ $\hat{b}_0$ 和 $\hat{b}_1$ ）；观测值数目和估计参数的个数之差为 $n - 2$ 。该差额也被称为自由度；这是保证所估计的估计量是无偏的所需要的

分母。

例 8-12 估计量标准误的计算

回忆一下对于图 8-8 中所给出的通货膨胀率和货币供应量增长率的数据，所估计的回归等式为 $Y_i = 0.0084 + 0.5545X_i$ 。表 8-7 使用该估计的等式来计算估计值标准误所需要的数据。

表 8-7 估计量标准误的计算

国家	货币供应量增长率 X_i	通货膨胀率 Y_i	预测的通货膨胀率 $\hat{Y}_i$	回归的残差 $Y_i - \hat{Y}_i$	残差的平方 $(Y_i - \hat{Y}_i)^2$
澳大利亚	0.1166	0.0676	0.0731	-0.0055	0.000030
加拿大	0.0915	0.0519	0.0591	-0.0072	0.000052
新西兰	0.1060	0.0815	0.0672	0.0143	0.000204
瑞士	0.0575	0.0339	0.0403	-0.0064	0.000041
英国	0.1258	0.0758	0.0782	-0.0024	0.000006
美国	0.0634	0.0509	0.0436	0.0073	0.000053
总和					0.000386

资料来源：国际货币基金。

表 8-7 的第 1 列和第 2 列的数字给出了 6 个国家的长期货币供应量增长率 X_i 和长期通货膨胀率 Y_i 。第 3 列的数字给出了拟合回归方程对于每个观测值的因变量的预测值。例如，对于美国来说，长期通货膨胀率的预测值为 $0.0084 + 0.5545(0.634) = 0.0436$ 或者 4.36%。倒数第 2 列包含了回归的残差，其是因变量 Y_i 的实际值和其预测值之差 ($\hat{Y}_i = \hat{b}_0 + \hat{b}_1 X_i$)。所以对于美国来说，残差等于 $0.0509 - 0.0436 = 0.0073$ 或者 0.73%。最后一列包含回归残差的平方。残差的平方和等于 0.000386。应用估计量标准误的公式，我们有

$$\left(\frac{0.000386}{6 - 2} \right)^{1/2} = 0.009823$$

因此，估计量的标准误约为 0.98%。

下面，我们将把该估计量和对于该回归中参数不确定性的估计结合起来，以确定从货币供应量增长率所预测到的通货膨胀率的置信区间。我们将看到标准误越小将使得预测越精确。

8.3.4 决定系数

虽然估计量的标准误给出一些关于我们对利用回归方差得到某个特定预测值 Y 确信程度的指标，但是它仍旧没有告诉我们自变量对于因变量变动的解释力有多强。而决定系数恰好完成了该任务：它度量了自变量对于因变量总变动的解释比率。

我们可以通过两种方式来计算决定系数。比较简单的方法（该方法可以用于单个自变量的线性回归）是将因变量和自变量的相关系数平方。例如，回忆一下，1970~2001 年 6 个工业化国家长期货币供应量增长率和长期通货膨胀率之间的相关系数为 0.8702。因此，图 8-8 所给出的回归的决定系数为 $(0.8702)^2 = 0.7572$ 。因此在这个回归中，长期货币供应量增长率解释了 1970~2001 年间不同国家的约 76% 的长期通货膨胀率的变动。

该方法所存在的问题是它不能够在具有一个以上自变量的情形下使用。[Ⓐ]因此，我们需要一种替代的方法来计算多元自变量的决定系数。现在我们给出该替代方法背后的逻辑思想。

如果我们不知道回归的关系，我们对于任何特定因变量观测值的猜测应该简单地为 $\bar{Y}$ ，因变量的均值。一个基于 $\bar{Y}$ 来度量 Y_i 预测的准确性的指标为 Y_i 的样本方差 $\sum_{i=1}^n \frac{(Y_i - \bar{Y})^2}{n - 1}$ 。一个替代的方法是利用 $\bar{Y}$ 来预测某个特定的观测值 Y_i ，该方法使用回归关系来进行预测。在这种情况下，我们的预测值将为 $\hat{Y}_i = \hat{b}_0 + \hat{b}_1 X_i$ 。如果回归关系拟合得很好，那么利用 $\bar{Y}$ 来预测 $\hat{Y}_i$ 的误差将远小于利用 $\bar{Y}$ 来预测 Y_i 的误差。如果我们将 $\sum_{i=1}^n (Y_i - \bar{Y})^2$ 称为 Y 的总变差，而将 $\sum_{i=1}^n (Y_i - \hat{Y}_i)^2$ 称为回归中未能解释的变差，那么我们可以利用下面的等式来度量回归所能解释的变差：

总变差 = 不能解释的变差 + 能被解释的变差 (8-7)

决定系数是由回归进行解释的总变差的一个分数。其给予我们如下关系：

$R^2 = \frac{\text{能被解释的变差}}{\text{总变差}} = \frac{\text{总变差} - \text{不能被解释的变差}}{\text{总变差}} = 1 - \frac{\text{不能被解释的变差}}{\text{总变差}}$ (8-8)

注意到总变差等于能被解释的变差加上不能被解释的变差，如同在式 (8-7) 中所给出的。大部分回归程序报告的决定系数为 R^2 。[Ⓑ]

例 8-13 通货膨胀率和货币供应量增长率

利用表 8-7 中的数据，我们可以看到由回归不能解释的变差（残差的平方和）等于 0.000 386。表 8-8 给出了因变量（长期通货膨胀率）总变差的计算。

表 8-8 总变差的计算

国家	货币供应量增长率 X_i	通货膨胀率 Y_i	偏离均值的偏差 $Y_i - \bar{Y}$	偏差的平方 $(Y_i - \bar{Y})^2$
澳大利亚	0.116 6	0.067 6	0.007 3	0.000 053
加拿大	0.091 5	0.051 9	-0.008 4	0.000 071
新西兰	0.106 0	0.081 5	0.021 2	0.000 449
瑞士	0.057 5	0.033 9	-0.026 4	0.000 697
英国	0.125 8	0.075 8	0.015 5	0.000 240
美国	0.063 4	0.050 9	-0.009 4	0.000 088
	平均值	0.060 3	总和	0.001 598

资料来源：国际货币基金。

该区间通货膨胀率的平均值为 6.03%。倒数第 2 列给出了每个国家长期通货膨胀率偏离其均值的偏差值；最后一列给出了该偏差值的平方。这些偏差的平方和为样本 Y 的总变差，如表 8-8 中所示。

该回归的决定系数为

$\frac{\text{总变差} - \text{不能被解释的变差}}{\text{总变差}} = \frac{0.001\,598 - 0.000\,386}{0.001\,598} = 0.758\,4$

注意到该方法给出与我们之前得到的（在最后的小数位上的差别是由于近似而造成的）相同结果（保留到小数点后两位）。我们将在多元回归一章中再次使用该方法；当我们具有多于一个自变量时，该方法是唯一能够计算出决定系数的方法。

Ⓐ 我们将在多元回归一章中讨论这种模型。
Ⓑ 正如我们在本章后面的回归输出表格中所展现的，回归程序同样也会给出复相关系数（multiple R ）的报告，其为 Y 的实际值和预测值之间的相关系数。该决定系数为多元 R 平方。

8.3.5 假设检验

在这一节中,我们将讨论关于一个回归模型中的截距和斜率总体值的假设检验问题。这个问题在实际应用中是非常重要的。例如,我们可能想要利用资本资产定价模型来检查一只股票的估值;我们假设该股票具有市场平均的贝塔或者系统风险水平。或者我们可能想要检验经济学家对于通货膨胀率的预测是否是无偏的(平均意义上,没有高估也没有低估)。在每个例子中,是否有证据能支持该假设?诸如这些问题可以通过回归模型的假设检验来回答。该检验经常是对于截距或者斜率的 t 检验。为了理解关于该检验的这些概念,回顾一下基于置信区间的简单和等价的方法对我们来说是十分有用的。

如果我们知道3件事,那我们就可以利用置信区间方法来进行假设检验:(1)估计参数的值 $\hat{b}_0$ 和 $\hat{b}_1$;(2)参数 b_0 和 b_1 的假定值;(3)包含估计参数的一个置信区间。一个置信区间是在一个给定置信水平下,我们认为包含真实参数值 b_1 的区间。为了计算出一个置信区间,我们必须为检验选择显著性水平,并知道估计参数的标准误。

假设我们用股票市场指数的收益率来对于单只股票的收益率进行回归,并求得斜率($\hat{b}_1$)为1.5,而且其标准误($s_{\hat{b}_1}$)为0.200。假设在我们的回归分析中使用62个月度观测值。参数(b_1)的假定值为1.0,这是市场的平均斜率。估计的总体斜率常被称为贝塔,因为总体的系数常用希腊符号贝塔(β),而不是我们在文中所用的字母 b_1 来表示。我们的原假设为 $b_1 = 1.0$ 以及 $\hat{b}_1$ 是 b_1 的估计量。我们对于我们的检验使用95%的置信区间,或者我们可以说该检验具有0.05的显著性水平。

我们的置信区间将跨越 $\hat{b}_1 - t_c s_{\hat{b}_1}$ 到 $\hat{b}_1 + t_c s_{\hat{b}_1}$ 的区间,或者

$$\hat{b}_1 \pm t_c s_{\hat{b}_1} \quad (8-9)$$

式中, t_c 表示 t 值的临界值。[⊖]该检验的临界值取决于原假设下 t 分布的自由度数值。自由度的数值等于观测数减去估计的参数数量。在单个自变量的回归模型中,有两个估计参数,它们分别为截距项和自变量前的系数。对于这个例子中的62个观测值和2个估计参数,我们的自由度为60(62-2)。对于自由度为60的 t 分布,书后临界值的数值表中显示,在0.05显著性水平下,该 t 值的临界值为2.00。将我们例子中的值代入式(8-9),我们得到区间

$$\begin{aligned} \hat{b}_1 \pm t_c s_{\hat{b}_1} &= 1.5 \pm 2.00 \times (0.200) \\ &= 1.5 \pm 0.400 \\ &= 1.10 \text{ 到 } 1.90 \end{aligned}$$

在原假设下,该置信区间包含 b_1 的概率为95%。因为我们检验的是 $b_1 = 1.0$,而该置信区间没有包含1.0,所以我们能够拒绝原假设。因此,我们有95%的信心认为该股票的贝塔不等于1.0。

在实际应用中,使用一个回归模型来检验假设的最常用方法是根据 t 检验的显著性。为了检验假设,我们可以计算统计量

$$t = \frac{\hat{b}_1 - b_1}{s_{\hat{b}_1}} \quad (8-10)$$

该检验统计量服从自由度为 $n-2$ 的 t 分布,因为在回归中有两个参数需要估计。我们将该 t 统计量的绝对值与 t_c 相比较。如果 t 的绝对值大于 t_c ,那么我们应该拒绝原假设。将上例中的数值代入该关系式,我们得到与股票贝塔等于1.0($b_1 = 1.0$)的概率相关的 t 统计量。

⊖ 我们对于该检验使用 t 分布,这是因为我们使用样本标准误的估计量 t_c ,而不是其真实(总体)的数值。在抽样和估计一章中,我们讨论了自由度的概念。

$$\begin{aligned}t &= \frac{\hat{b}_1 - b_1}{s_{\hat{b}_1}} \\&= \frac{1.5 - 1.0}{0.200} \\&= 2.50\end{aligned}$$

因为 $t > t_c$, 所以我们拒绝 $b_1 = 1.0$ 的原假设。

上例中的 t 统计量为 2.50, 并且在 0.05 显著性水平下, $t_c = 2.00$; 因此我们拒绝了原假设, 因为 $t > t_c$ 。该陈述等价于说我们有 95% 的信心认为关于斜率的区间没有包含值 1.0。然而, 如果我们在 0.01 显著性水平下进行该检验, t_c 将为 2.66, 而我们将无法拒绝原假设, 因为 t 值在该显著性水平下没有大于 t_c 。关于斜率的 99% 的置信区间将包含数值 1.0。

显著性水平的选择总是一个值得考虑的问题。当我们使用更高的置信水平, t_c 将增加。该选择将导致更宽的置信区间, 从而降低拒绝原假设的可能性。分析师们常常选择 0.05 的显著性水平, 其表示当事实上该假设为真的情况下, 有 5% 的机会来拒绝原假设 (第一类错误)。当然, 将显著性水平从 0.05 降低到 0.01 会使得第一类错误发生的概率下降, 但是同时它会增加第二类错误发生的可能性——当事实上原假设为假时, 没有拒绝原假设。[⊖]

通常, 金融分析师们并不会简单地报告他们的假设是否会拒绝某个关于回归参数的特定假设。相反, 他们会报告关于某个特定假设的 p 值或者概率值。 p 值是指能拒绝原假设的最小的显著性水平。它使得读者能够自己解释结果, 而不是要被告知某个特定的假设已经被拒绝或者被接受。在大部分的回归软件包中, 给定估计的系数和系数的标准误, 对于原假设 (真实的参数等于 0) 对应于备择假设 (该参数不等于 0) 的检验, 会使用所给出的关于回归系数的 p 值。例如, 如果 p 值为 0.005, 那么我们可以在 0.5% 的显著性水平下 (99.5% 的置信度) 拒绝真实的参数等于 0 的原假设。

估计系数的标准误是关于回归系数假设检验 (和估计系数的置信区间) 的重要输入变量。较强的回归结果将导致估计参数的标准误较小, 并得到较窄的置信区间。如果上例中的标准误 ($s_{\hat{b}_1}$) 为 0.100, 而不是 0.200, 该置信区间的范围将是原来的一半, 而 t 统计量将变为原来的两倍。在标准误如此小的情况下, 我们甚至在 0.01 的显著性水平下也将拒绝原假设, 因为我们有 $t = (1.5 - 1)/0.1 = 5.00$, 而 $t_c = 2.66$ 。

在此背景下, 我们可以对于利用实际回归结果的假设检验进行讨论。下面 3 个例子举例说明了在不同的典型投资环境中假设检验的应用。

例 8-14 通用汽车股票贝塔的估计

你是一名投资于通用汽车股票的投资者, 而你想要获得该股票贝塔的估计值。如同在文中所举的例子, 你假定 GM 具有市场平均的风险水平, 并且它超过无风险收益率的要求回报率与市场所要求的超额收益率一致。描述这些表述的一个回归方程为

$$(R - R_F) = \alpha + \beta(R_M - R_F) + \varepsilon \quad (8-11)$$

式中, R_F 表示区间的无风险收益率 (在期初就已知的), R_M 表示区间的市场收益率, R 表示区间公司股票的收益率, 而 β 等于公司收益率与市场收益率的协方差除以市场收益率的方差, $\text{Cov}(R, R_M)/\sigma_M^2$ 。对该等式用线性回归来估计, 得到 β 的估计值 $\hat{\beta}$, 该值告诉我们在给定市场

⊖ 对于第一类错误和第二类错误的充分讨论, 参见假设检验一章。

收益率预期的条件下，该证券所要求的收益率溢价的大小为多少。[Ⓐ]

假设我们想要检验原假设 H_0 : GM 股票的 $\beta = 1$ ，来观察 GM 股票是否在总体上具有和市场所要求的收益率一样的溢价。我们需要 GM 股票的收益率、无风险利率和市场指数收益率的数据。对于这个例子，我们使用从 1998 年 1 月到 2002 年 12 月的数据 ($n = 60$)。GM 股票的收益率为 R 。30 天国库券的月度收益率为 R_F 。标准普尔 500 的收益率为 R_M 。[Ⓑ]我们对于两个参数进行估计，因此自由度的数值为 $n - 2 = 60 - 2 = 58$ 。表 8-9 给出了回归 $(R - R_F) = \alpha + \beta(R_M - R_F) + \varepsilon$ 的结果。

表 8-9 对于 GM 股票贝塔的估计

回归统计结果			
复相关系数	0.554 9		
R 平方	0.307 9		
估计量的标准误	0.098 5		
观测数目	60		
	系数	标准误	t 统计量
阿尔法	0.003 6	0.012 7	0.284 0
贝塔	1.195 8	0.235 4	5.079 5

资料来源：伊博森协会和彭博 L.P.

我们要检验原假设 H_0 : GM 的 β 等于 1 ($\beta = 1$) 对应于备择假设: β 不等于 1 ($\beta \neq 1$)。从回归中所得到的估计值 $\hat{\beta}$ 为 1.195 8。回归中该系数的标准误 $s_{\hat{\beta}}$ 为 0.235 4。回归方程的自由度为 58 ($60 - 2$)，所以在 0.05 显著性水平下，该检验统计量的临界值近似为 $t_c = 2.00$ 。因此，对于任一 β 的假定值，该数据所给出的 95% 置信区间为

$$\begin{aligned} &\hat{\beta} \pm t_c s_{\hat{\beta}} \\ &1.195 8 \pm 2.00 \times (0.235 4) \\ &0.725 0 \text{ 到 } 1.666 6 \end{aligned}$$

在这个例子中，假定参数值为 $\beta = 1$ ，而数值 1 将落入该置信区间之中，所以我们在 0.05 显著性水平下，无法拒绝原假设。这意味着我们无法拒绝 GM 股票在总体上具有和市场相同系统风险的原假设。

对于该问题的另一种考虑方式为利用式 (8-10) 计算 GM 贝塔假定参数的 t 统计量：

$$t = \frac{\hat{\beta} - \beta}{s_{\hat{\beta}}} = \frac{1.195 8 - 1.0}{0.235 4} = 0.831 8$$

这个 t 值要比临界 t 值 2.00 要小。所以，这两种方法都没有导致我们拒绝原假设。注意到在表 8-9 的回归结果中与 $\hat{\beta}$ 相关的 t 值为 5.079 5。给定我们使用的显著性水平，我们无法拒绝 $\beta = 1$ 的原假设，但是我们可以拒绝 $\beta = 0$ 的原假设。[Ⓒ]

Ⓐ 贝塔 (β) 一般会用 60 个月的历史数据进行估计，但是数据样本长度有时会有些变动。虽然月度数据常常被使用，但是一些金融分析师们也会使用日数据来估计 β 。关于估计 β 方法的更多内容，参见 Reilly 和 Brown (2003)。给定某个高于无风险收益率的市场超额收益率 ($R_M - R_F$)，高于无风险收益率的 GM 股票的期望超额收益率 ($R - R_F$) 为 $\beta(R_M - R_F)$ 。该结果成立的原因是我们将 $(R - R_F)$ 与 $(R_M - R_F)$ 进行回归。例如，如果一只股票的贝塔为 1.5，那么其期望超额收益率为市场投资组合超额收益率的 1.5 倍。

Ⓑ GM 股票收益率的数据来源于彭博数据库。国库券收益率和标准普尔 500 的收益率数据来自伊博森协会。

Ⓒ 一个系数的 t 统计量会自动由统计软件程序报告出来，该值假设的原假设为其系数为 0。如果你需要一个不同的原假设，正如我们在这个例子中所做的 ($\beta = 1$)，你必须自行构建正确的检验统计量或者指示程序来计算它。

同时注意到这个回归的 R^2 仅为 0.3079。该结果表明只有约 31% 的 GM 股票超额收益率（高于无风险收益率的 GM 收益率）的总变差可以被市场组合的超额收益率所解释。GM 股票超额收益率总变差的剩余 69% 是非系统部分，其可以被归咎于公司特定的风险。

在下面的例子中，我们给出单边备择假设的回归假设检验。

例 8-15 利用相对于投资资本的收益率来解释公司价值

一些金融分析师们认为度量公司创造财富能力的最好方法为将公司投资资本的收益率（ROIC）和加权平均资本成本（WACC）相比较。如果一家公司的 ROIC 小于其资本成本，就说明它正在侵蚀财富。[⊖]

企业价值（EV）为基于市场价格的度量公司价值的指标，其被定义为股票和债券的市场价值减去现金和投资的价值。投资资本（IC）为公司价值的会计度量，其被定义为股票和债券的账面价值之和。EV 相对于 IC 的比率越高，一般说明创造财富的成功可能性越大。茂宝欣（Mauboussin）（1996）认为 ROIC 和 WACC 之间的差额解释了 EV 相对于 IC 的比率。使用食品制造业公司的数据，我们可以利用式（8-12）中所给出的回归模型，检验 EV/IC 和（ROIC-WACC）之间的关系。

$$\frac{EV_i}{IC_i} = b_0 + b_1 (ROIC_i - WACC_i) + \varepsilon_i \tag{8-12}$$

式中，下标 i 是识别不同公司的指标。我们的原假设为 $H_0: b_1 \leq 0$ ，我们设显著性水平为 0.05。如果我们拒绝原假设，那么我们就具有统计上显著的关于 EV/IC 和（ROIC-WACC）之间相关关系的证据。我们利用 2001 年 9 家食品加工业公司的数据来估计式（8-12）。[⊖]表 8-10 和图 8-11 给出了该回归的结果。

表 8-10 利用 ROIC-WACC 的差额来解释企业价值/投资资本

回归的统计结果			
复相关系数	0.9469		
R 平方	0.8966		
估计量标准误	0.7422		
观测数目	9		
	系数	标准误	t 统计量
截距	1.3478	0.3511	3.8391
差额	30.0169	3.8519	7.7928

资料来源：尼尔森（2003）。

我们基于估计斜率系数的 t 统计量约为 7.79 的结果，拒绝原假设。在我们这些公司的样本中收益率差额（ROIC-WACC）和 EV/IC 的比率之间的存在较强的正向关系。图 8-11 展现出了这一较强的正向关系。0.8966 的 R^2 表明收益率的差额解释了 90% 的 2001 年样本中食品加工公司的 EV/IC 比率的变差。对我们这些公司的样本来说，收益率差额前的系数 30.0169 说明对于收益率差额每一个百分点的增加，预测的 EV/IC 的增加为 $0.01(30.0169) = 0.3002$ 或者 30%。

⊖ 例如，参见 Stewart（1991）和 Mauboussin（1996）。
⊖ 我们的数据来源于 Nelson（2003）。许多销售端分析师使用这类回归。该类回归是在出版的分析师报告中最常使用的横截面回归之一。

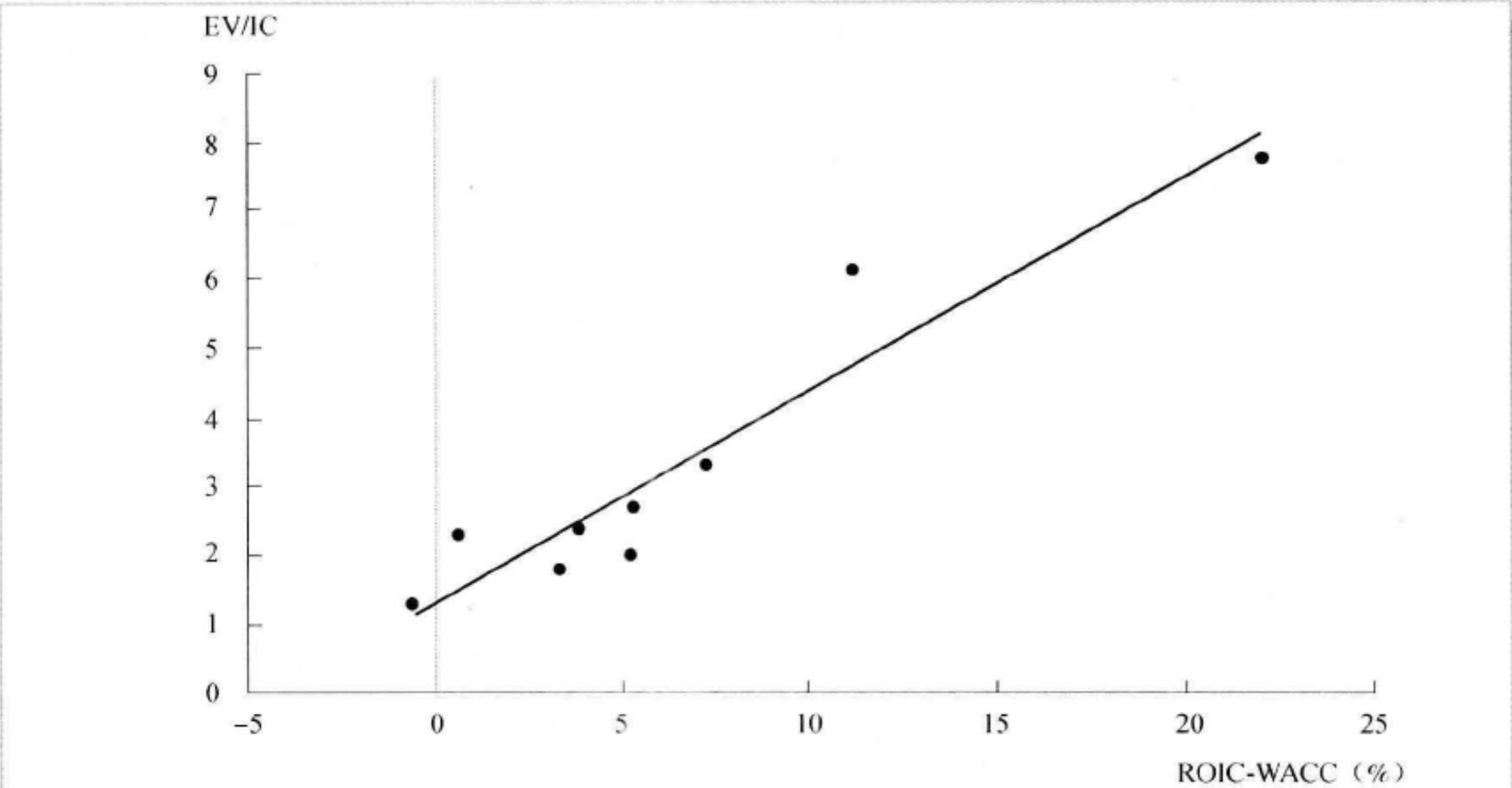


图 8-11 利用食品行业的 ROIC-WACC 差额来解释企业价值/投资资本的拟合回归直线
资料来源：CSFB 食品投资者手册 2003。

在本节最后的这个例子中，我们将给出涉及与斜率为 0 相同的斜率为 1 的原假设。

例 8-16 检验是否通货膨胀率的预测是无偏的

例 8-11 介绍了检验预测偏差的概念。那个例子表明如果一个预测值是无偏的，那么其预期误差为 0。我们能通过将每个时点上的预测值与在预测值公布之后经济变量的实际值相比较来检验是否对于特定经济变量时间序列的预测值是无偏的。如果预测值是无偏的，那么根据定义，已实现预测值的平均误差应该接近于 0。在那个例子中， b_0 （截距）的值应该为 0，而且 b_1 （斜率）的值应该为 1，如同例 8-11 中所讨论的。

再一次查看图 8-9，该图给出了 1983 年第一季度到 2002 年第四季度（ $n=80$ ）由职业经济预测家给出的对于 CPI 百分比变动的当季度预测值以及实际百分比的变动。为了检验预测值是否是无偏的，我们必须估计例 8-11 中给出的回归方程。我们在表 8-11 中报告了该回归的结果。所估计的等式为

$$\text{CPI}_t \text{ 的实际百分比变动} = b_0 + b_1 (\text{预测的变动}_t) + \varepsilon_t$$

这一回归对于两个参数（截距和斜率）进行估计；因此，该回归的自由度为 $n - 2 = 80 - 2 = 78$ 。

表 8-11 检验是否 CPI 的预测值是无偏的
(因变量：以百分比来表述的 CPI 变动)

回归的统计结果			
复相关系数	0.713 8		
R^2	0.509 5		
估计量的标准误	1.032 2		
观测数目	80		
	系数	标准误	t 统计量
截距	-0.014 0	0.365 7	-0.038 4
预测（斜率）	0.963 7	0.107 1	9.000 8

资料来源：费城和圣路易斯联邦储备银行。

我们现在能够对于该回归中两个参数的原假设进行检验了。我们的第1个原假设为该回归的截距为0($H_0: b_0 = 0$)。备择假设为截距不等于0($H_a: b_0 \neq 0$)。我们第2个原假设为该回归的斜率等于1($H_0: b_1 = 1$)。其备择假设为斜率不等于1($H_a: b_1 \neq 1$)。

为了检验关于 b_0 和 b_1 的假设,我们必须首先确定一个基于特定显著性水平的临界值以及构建对于每个参数的置信区间。如果我们选择0.05的显著性水平,在自由度为78的情况下,临界值 t_c 约为1.99。参数 $\hat{b}_0$ 的估计值为-0.0140, $\hat{b}_0$ 标准误的估计值($s_{\hat{b}_0}$)为0.3657。令 B_0 代表任一假定值。因此,在 $b_0 = B_0$ 的原假设下, b_0 的一个95%置信区间为

$$\begin{aligned} & \hat{b}_0 \pm t_c s_{\hat{b}_0} \\ & -0.0140 \pm 1.99 \times (0.3657) \\ & -0.7417 \text{ 到 } 0.7137 \end{aligned}$$

在这个例子中, B_0 为0。由于0这个值落入置信区间之中,所以我们无法拒绝第1个 $b_0 = 0$ 的原假设。我们将在之后马上说明如何来解释这一结果。

我们第2个原假设是基于与我们第1个原假设的检验中使用的相同样本。因此,检验该假设的临界值和检验第1个原假设的临界值($t_c = 1.99$)是一样的。参数 $\hat{b}_1$ 的估计值为0.9637,而 $\hat{b}_1$ 标准误的估计值 $s_{\hat{b}_1}$ 为0.1071。因此,对于 b_1 的任一特定假定值的95%置信区间可以构造如下:

$$\begin{aligned} & \hat{b}_1 \pm t_c s_{\hat{b}_1} \\ & 0.9637 \pm 1.99 \times (0.1071) \\ & 0.7506 \text{ 到 } 1.1768 \end{aligned}$$

在这个例子中,我们对于 b_1 的假定值为1。由于数值1落入置信区间之中,所以在0.05显著性水平下,我们无法拒绝 $b_1 = 1$ 的原假设。因为我们不能拒绝关于这个模型中参数的两个原假设中的任何一个($b_0 = 0, b_1 = 1$),所以我们无法拒绝对于CPI变动的预测值是无偏的假设。^①

作为一名分析师,你经常需要对于经济增长率的预测来帮助你做出关于资产配置、预期收益率以及其他投资决策的推荐。刚才所进行的假设检验表明你无法拒绝在职业预测家调查中的CPI预测值是无偏的假设。如果你为了你的资产配置决策,需要一个关于未来CPI百分比变动的无偏预测值,你可能会使用这些预测值。

8.3.6 单变量回归中的方差分析

方差分析(analysis of variance, ANOVA)是将一个变量的总方差分解成归因于不同来源部分的一种统计方法。^②在回归分析中,我们使用ANOVA来确定一个或者多个自变量在解释因变量变差中的有用性。在方差分析中所采取的一个重要的统计检验是F检验。F统计量检验所有在线性回归中的斜率是否都等于0。在单个变量的回归中,这个检验的原假设为 $H_0: b_1 = 0$ 对应于备择假设 $H_a: b_1 \neq 0$ 。

为了准确地确定对于斜率等于0的原假设的检验统计量,我们需要了解下列变量:

- ① 对于假设 $b_0 = 0$ 和 $b_1 = 1$ 的联合检验将要求我们把 $\hat{b}_0$ 和 $\hat{b}_1$ 的协方差加以考虑。关于这类联合假设检验的内容,参见Greene(2003)。
- ② 在本章中,我们关于ANOVA在回归中的应用,这是金融分析师们会使用该工具的最常见的情形。在这一情形下,ANOVA被用来检验所有的回归斜率是否都等于0。分析师们也使用ANOVA来检验两个或者多个总体的均值都相等的假设。更多内容详见Daniel和Terrell(1995)。

- 观测值的总数目 (n);
- 所要估计参数的总数目 (在单个自变量的回归中, 该数目为两个: 截距和斜率);
- 误差或者残差的平方和 (the sum of squared errors or residuals), $\sum_{i=1}^n (Y_i - \hat{Y}_i)^2$, 简称为 SSE。该值也被称为剩余 (残余) 平方和 (the residual sum of squares);
- 回归平方和 (the regression sum of squares), $\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2$, 简称为 RSS。该值是回归方程中能被解释的 Y 的总变差的数值。总变差 (total variation, TSS) 是 SSE 和 RSS 的总和。

确定斜率是否等于 0 的 F 检验基于的是由这 4 个值所构建的 F 统计量。 F 统计量度量的是回归等式解释因变量变差的好坏。 F 统计量是平均的回归平方和与平均的误差平方和之比。平均回归平方和是通过将回归平方和除以所估计斜率参数的数目 (在这个例子中为一个) 而计算得到的。平均误差平方和是通过误差平方和除以观测值的数目 n 减去所估计参数的总数 (在这个例子中为两个: 截距和斜率) 而计算得到的。这两个除数是 F 检验的自由度。如果有 n 个观测值, 对于斜率等于 0 的原假设的 F 检验, 在这里记为 $F_{\text{#斜率参数}, n-\text{#参数}} = F_{1, n-2}$, 而该检验的自由度为 1 和 $n-2$ 。

例如, 假设回归模型中的自变量不能解释任何因变量的变动。那么对于回归模型的预测值 $\hat{Y}_i$ 是因变量的平均值 $\bar{Y}$ 。在这个例子中, 回归平方和 $\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2$ 为 0。因此, F 统计量为 0。如果自变量只能解释因变量变动的很小一部分, 那么 F 统计量的值也将会很小。

在单个自变量回归中 F 统计量的公式为:

$$F = \frac{\text{RSS}/1}{\text{SSE}/(n-2)} = \frac{\text{平均回归平方和}}{\text{平均误差平方和}} \quad (8-13)$$

如果回归模型能很好地解释因变量的变动, 那么该比率应该会较高。每单位所估计参数所解释的回归平方和相对于每单位自由度不能解释的变动将会较高。在本书书后给出了这一 F 统计量临界值的数值表。

即使 F 统计量常通过回归软件包来进行计算, 分析师们一般也不在只有单个自变量的回归模型中使用 ANOVA 和 F 检验。为什么不用呢? 在这种回归中, F 统计量是关于斜率 t 统计量的平方。因此, F 检验重复了检验斜率显著性的 t 检验。这一关系在两个或者多个斜率系数的回归模型中是不成立的。虽然如此, 单个斜率系数的例子给予我们理解多个斜率系数例子的基础。

通常, 共同基金业绩的表现基于基金是否具有正的阿尔法——显著的正的超额风险调整后的收益率[⊖]来进行评估的。一个常用的风险调整方法是基于资本资产定价模型。考虑这个回归

$$(R_i - R_F) = \alpha_i + \beta_i(R_M - R_F) + \varepsilon_i \quad (8-14)$$

式中, R_F 表示区间无风险收益率 (在期初就已经知道的), R_M 表示区间市场的收益率, R_i 表示区间共同基金 i 的收益率, 而 β_i 表示基金的贝塔。如果 $\alpha_i = 0$, 那么该基金具有的风险调整后的超额收益率为 0。如果 $\alpha_i > 0$, 那么 $(R_i - R_F) = \beta_i(R_M - R_F) + \varepsilon_i$, 而取期望之后, $E(R_i) = R_F + \beta_i(R_M - R_F)$, 这表明 β_i 完全解释了基金平均的超额收益率。例如, 如果 $\alpha_i > 0$, 那么基金将获得比给定贝塔的预期要更高的收益率。

总之, 为了检验一个基金是否具有一个正的阿尔法, 我们必须检验基金不具有风险调整后的超

⊖ 注意到希腊字母阿尔法 α 在传统上是用来表示式 (8-14) 中的截距, 而它不应该与另一个将 α 用来表示显著性水平的传统用法相混淆。

额收益率的原假设 ($H_0: \alpha = 0$) 对应于风险调整后的超额收益率不为 0 的备择假设 ($H_a: \alpha \neq 0$)。

例 8-17 业绩评价：德雷福斯 (Dreyfus) 增值基金

表 8-12 给出了 1998 年 1 月到 2002 年 12 月德雷福斯 (Dreyfus) 增值基金超额收益率的评估结果。注意到在回归中所估计的贝塔 $\hat{\beta}_1$ 为 0.790 2。德雷福斯增值基金在总体上被认为具有约 0.8 倍的市场风险。

表 8-12 1998 年 1 月 ~ 2002 年 12 月德雷福斯增值基金的业绩评价

回归的统计结果				
复相关系数		0.928 0		
R 平方		0.861 1		
估计量的标准误		0.017 4		
观测数目		60		
ANOVA	自由度 (df)	平方和 (SS)	平均平方和 (MSS)	F
回归	1	0.109 3	0.109 3	359.64
残差	58	0.017 6	0.000 3	
总数	59	0.126 9		
	系数	标准误	t 统计量	
阿尔法	0.000 9	0.002 3	0.403 6	
贝塔	0.790 2	0.041 7	18.965 5	

资料来源：芝加哥大学证券价格研究中心。

同时注意到该回归中所估计的阿尔法 ($\hat{\alpha}$) 为正数 (0.000 9)。系数的值仅仅比该系数的标准误的 1/3 多一点点 (0.002 3)，所以该系数的 t 统计量仅为 0.403 6。因此，我们无法拒绝基金不具有显著高于与基金的市场风险相关的收益率的超额收益率的原假设 ($\alpha = 0$)。该结果表明基金的收益率可以由基金的市场风险来解释，而且在这个区间中，对于该基金的超额收益率没有表现出额外的统计上的显著性。[⊖]

因为该回归中斜率系数的 t 统计量为 18.965 5，该系数的 p 值小于 0.000 1，近似于 0。因此，这一系数真实值实际上为 0 的概率是微小的。

那么我们应该如何使用 F 检验来确定这一回归的斜率是否等于 0？表 8-12 中 ANOVA 的这部分提供了我们所需的数据。在这一例子中，

- 观测值的总数目 (n) 为 60；
- 所要估计参数的总数目为 2 个 (截距和斜率)；
- 误差或者残差的平方和 SSE 为 0.017 6；
- 回归平方和 RSS 为 0.109 3。

因此，检验是否斜率系数为 0 的 F 统计量为

$$\frac{0.109\,3/1}{0.017\,6/(60-2)} = 360.19$$

⊖ 这个例子介绍了一个广为人知的利用资本资产定价模型回归的投资方法。然而，研究者们发现了解释由线性回归所获得阿尔法中的一些问题。所管理的投资组合的系统性风险是受到投资组合经理控制的。如果投资组合的贝塔与市场收益率因此是相关的 (可能是由于市场择时所造成的)，那么基于最小二乘法的贝塔所推断的阿尔法，如同这里所获得，可能会产生错误。这一高级的课题在 Dybvig 和 Ross (1985a) 和 (1985b) 中进行了讨论。

(与表 8-12 中 F 统计量的细微差别源于四舍五入。) ANOVA 的输出结果给出该 F 统计量的 p 值小于 0.0001, 而该值正好与关于斜率的 t 统计量的 p 值相同。因此, F 检验并没有告诉比我们已从 t 检验中所知道的更多信息。我们同样也能看到 F 统计量 (359.64) 是 t 统计量 (18.9655) 的平方。

8.3.7 预测区间

金融分析师们经常想要利用回归结果来对于因变量做出预测。例如, 我们可能会问, “如果实际 GDP 以 4% 的速度增长, 那么今年 XYZ 公司的销售额会增长多快?” 但是我们不仅仅对于做出这些预测感兴趣; 我们也想要知道应该对未来结果抱有多大的信心。例如, 如果我们预测今年 XYZ 公司的销售额会增长 6%, 那么我们的预测将意味着更大的发生该情况的可能性, 如果我们以 95% 的信心认为销售额增长将落入 5% ~ 7% 的区间, 而不是只有 25% 的信心认为这个结果会发生。因此, 我们需要理解如何计算围绕着回归预测值的置信区间。

当利用回归模型 $Y_i = b_0 + b_1 X_i + \varepsilon_i$, $i = 1, 2, \dots, n$ 和估计参数 $\hat{b}_0$ 和 $\hat{b}_1$ 来进行预测时, 我们必须考虑不确定性的两个来源。首先, 误差项本身包含着不确定性。误差项 σ_ε 的标准差可以用回归方程中估计值的标准误进行估计。然而, 第 2 个对于 Y 进行预测的不确定性的来源来自估计参数 $\hat{b}_0$ 和 $\hat{b}_1$ 的不确定性。

如果我们知道了回归参数 $\hat{b}_0$ 和 $\hat{b}_1$ 的真实值, 那么给定任一特定 X 的预测 (或假定) 值, 我们对于 Y 预测的方差将简单地记为 s^2 , 估计值标准误的平方。因为预测值 $\hat{Y}$ 将来自等式 $\hat{Y} = b_0 + b_1 X$ 和 $(Y - \hat{Y}) = \varepsilon$, 所以方差将为 s^2 。

然而, 因为我们必须估计回归参数 $\hat{b}_0$ 和 $\hat{b}_1$, 在给定任一 X 的预测值的情况下, 我们对于 Y 的预测值 $\hat{Y}$ 事实上应该为 $\hat{Y} = \hat{b}_0 + \hat{b}_1 X$ 。给定 X , Y 预测误差的估计方差 s_f^2 为

$$s_f^2 = s^2 \left[1 + \frac{1}{n} + \frac{(X - \bar{X})^2}{(n-1)s_x^2} \right] \quad (8-15)$$

估计方差取决于:

- 估计值标准误的平方 s^2 ;
- 观测值的数目 n ;
- 用来预测因变量的自变量 X 的值;
- 估计均值 $\bar{X}$;
- 自变量的方差 s_x^2 。[⊖]

一旦我们有了预测误差方差的估计值, 那么确定一个围绕预测值的预测区间就与本章之前所介绍的估计一个围绕估计参数的置信区间相类似。我们需要采取下列 4 个步骤来确定预测值的预测区间:

1. 作出预测。
2. 利用式 (8-15) 计算预测误差的方差。
3. 对于预测, 选择一个显著性水平 α 。例如, 给定回归中的自由度, 0.05 显著性水平确定了预测区间的临界值 t_c 。
4. 计算对于预测值的 $(1 - \alpha)\%$ 预测区间, 即 $\hat{Y} \pm t_c s_f$ 。

⊖ 对于该等式的推导, 参见 Pindyck 和 Rubinfeld (1998)。

例 8-18 预测公司价值相对于投资资本的比率

我们通过收益率相对于投资资本的比率和加权平均资本成本（ROIC-WACC）的差额来继续考察对于食品制造公司的公司价值相对于投资资本比率进行解释的例子。对于例 8-15，我们在表 8-10 中给出了估计的回归结果。

表 8-10（续） 通过 ROIC-WACC 的差额来解释公司价值/投资资本

回归统计量			
多元回归的 R 值	0.9469		
R 平方	0.8966		
估计值的标准误	0.7422		
观测值	9		
	系数	标准误	t 统计量
截距	1.3478	0.3511	3.8391
差额	30.0169	3.8519	7.7928

资料来源：尼尔森（2003）。

如果 ROIC 和 WACC 之间的差额收益率为 10%，那么你关心对于该公司的公司价值相对于投资资本比率的预测。

对于该公司的公司价值相对于投资资本比率的 95% 的置信区间为多少？

利用表 8-10 中所提供的数据，实施下列步骤：

1. 作出预测：期望 $EV/IC = 1.3478 + 30.0169(0.10) = 4.3495$ 。该回归表明如果 ROIC 和 WACC 的差额收益率（ X_i ）为 10%，那么 EV/IC 比率将为 4.3495。

2. 计算预测误差的方差。为了计算预测误差的方差，我们必须知道等式中估计值的标准误， $s = 0.7422$ （如同表 8-10 中所显示的）；

平均收益率差额 $\bar{X} = 0.0647$ （这个计算没有在表中给出）；以及样本中平均收益率差额的方差 $s_x^2 = 0.004641$ （这一计算没有在表中给出）。

利用这些数据，你可以计算在 10% 的 ROIC 和 WACC 间差额的情况下，对于该公司预测 EV/IC 的预测误差（ s_f^2 ）的方差。

$$\begin{aligned} s_f^2 &= 0.7422^2 \left[1 + \frac{1}{9} + \frac{(0.10 - 0.0647)^2}{(9 - 1)0.004641} \right] \\ &= 0.630556 \end{aligned}$$

在这个例子中，预测误差的方差为 0.630556，以及预测误差的标准差为 $s_f = (0.630556)^{1/2} = 0.7941$ 。

3. 确定 t 统计量的临界值。给定一个 95% 的置信区间以及 $9 - 2 = 7$ 的自由度，利用书后的数值表，t 统计量的临界值 t_c 为 2.365。

4. 计算预测区间。EV/IC 的 95% 置信区间是从 $4.3495 - 2.365(0.7941)$ 到 $4.3495 + 2.365(0.7941)$ ，或者 2.4715 ~ 6.2275。

总之，如果 ROIC 和 WACC 之间的差额为 10% 的话，那么 EV/IC 的 95% 的预测区间将从 2.4715 ~ 6.2275。较小的样本量被反映在相对较大的预测区间中。

8.3.8 回归分析的局限

虽然本章介绍了许多金融分析中所使用的回归模型，但是回归模型确实存在一些局限性。首先，回归关系可能随着时间发生改变，正如相关系数会发生变化一样。这个事实被称为参数的不稳定性（parameter instability）问题，而且这种不稳定性的存在性也不应该令人惊讶，因为运行的金融市场中的经济、税收、政府管制、政治以及机构的情况都会发生改变。无论考虑横截面回归还是时间序列回归，分析师都可能遇到这个问题。作为一个例子，股票特征间的横截面回归关系可能在成长型占主导的市场和价值型占主导的市场中表现不同。作为另一个例子，估计贝塔的时间序列回归经常会根据不同时间段的选取产生显著不同的估计贝塔。在横截面和时间序列这两种情况下，最常见的问题是从多于一个总体中抽样，当我们在这样做的时候如何对其进行识别也是一个具有挑战性的问题。

使用于特定投资情况下的回归结果的第2个局限性为回归关系的公共信息可能使得其对于未来的数据失效。例如，假设一位分析师发现具有某个特征的股票具有历史上非常高的收益率。如果其他的分析师也发现并利用了这一关系，那么具有该特征股票的价格将会被抬高。这一关系的信息可能导致该关系在未来不再存在。

最后，如果3.2节中所列出的回归假设被违背，那么基于线性回归的假设检验和预测将不再有效。虽然对于违背回归假设的情况有许多其他检验方法，但是对于一个假设是否被违背经常存在着许多不确定性。这一局限性将在多元回归一章中加以详细的讨论。

多元回归和回归分析中的问题

9.1 引言

作为一位金融分析师，我们经常需要使用一些比关于单个自变量的相关分析和回归更加复杂的统计方法。例如，一个对于纳斯达克股票的交易成本感兴趣的交易台可能需要了解确定纳斯达克买卖价差因素的信息，而一位共同基金的分析师可能想要了解一家投资于高科技股票共同基金的收益率表现是怎样的（是像一个成长型股票指数的收益率，还是更像一个价值型股票指数的收益率）。此外，一个投资者可能对于分析师推荐某只股票的决定因素感兴趣。对于这些问题，我们都能够利用一个关于多个自变量的线性回归——多元线性回归来进行回答。

在本章第2节和第3节中，我们将介绍并举例说明多元回归分析中的一些基本概念和模型。这些模型依赖于一些有时在实践中可能被违背的假设条件。在第4节中，我们将讨论3种主要的违背回归假设的情况。我们会重点介绍如何诊断假设是否被违背，以及当模型假设不满足时，如何采取相应的补救方法等实际问题。第5节归纳了一些建立良好回归模型所要遵循的准则，同时我们也讨论了一些分析师在建模过程中可能会犯的错误。在许多投资应用中，我们会对两个结果之一是否发生的概率感兴趣：例如，我们可能对于一只股票是否会被分析师推荐感兴趣。第6节将讨论一类能够回答上述问题的模型，这类模型被称为因变量为定性变量的模型。

9.2 多元线性回归

作为一位投资分析师，我们需要经常假设我们所关注变量的变动可以由多个变量来进行解释。我们所希望解释的变量被称为因变量。而我们认为可以用来解释因变量的这些变量被称为自变量。[⊖]一个能够使我们检验这两种变量间关系（如果它们之间存在关系的话）的工具就是多元线性回归。多元线性回归（Multiple linear regression）使得我们能够确定多个自变量对于某个特定因变量的影响情况。

⊖ 自变量也被称为解释变量或者回归量。

下面我们举一个例子来具体说明如何使用这一工具。假如我们想要了解交易商市场中所交易的股票的买卖价差是否受到交易该股票的做市商（交易商）数量和该股票市值的影响。我们可以利用下面的多元线性回归模型（multiple linear regression model）来对于该问题进行解答：

$$Y_i = b_0 + b_1 X_{1i} + b_2 X_{2i} + \epsilon_i$$

式中， Y_i 表示第 i 只股票买卖价差的自然对数（因变量）； X_{1i} 表示第 i 只股票做市商数量的自然对数； X_{2i} 表示第 i 家公司市场总价值的自然对数； ϵ_i 表示误差项；

当然，线性回归模型可以使用多于两个自变量来对于因变量进行解释。一个多元线性模型的一般形式为：

$$Y_i = b_0 + b_1 X_{1i} + b_2 X_{2i} + \cdots + b_k X_{ki} + \epsilon_i, \quad i = 1, 2, \cdots, n \quad (9-1)$$

式中， Y_i 表示第 i 个因变量 Y 的观测值； X_{ji} 表示第 i 个自变量 X_j 的观测值， $j = 1, 2, \cdots, k$ ； b_0 表示该回归方程的截距； $b_1, \cdots, b_k$ 表示每个自变量的斜率系数； ϵ_i 表示误差项； n 表示总的观测次数；

斜率系数 b_j 度量的是在其他自变量保持不变的情况下，自变量 X_j 变化一单位所造成的因变量 Y 的变化大小。例如，如果 $b_1 = 1$ ，而其他的自变量保持不变，那么我们可以预计 X_1 每增加一单位， Y 也将同样增加一单位。如果 $b_1 = -1$ ，而其他的自变量保持不变，那么我们可以预计 X_1 每增加一单位， Y 将减少一单位。多元线性回归可以估计出 $b_0, \cdots, b_k$ 的值。在本章中，我们将截距 b_0 和斜率系数 $b_1, \cdots, b_k$ 统称为回归系数（regression coefficients）。在我们接下去的讨论中，读者需要记住一个回归方程有 k 个斜率系数和 $k+1$ 个回归系数。

在实践中，我们常使用软件来对一个多元回归模型进行估计。例 9-1 给出了在实际投资中对于多元回归分析的一个应用。在讨论假设检验的过程中，例 9-1 给出了多元线性回归的典型回归结果以及其各项所对应的解释。

■ 例 9-1 买卖价差的解释

作为一家投资管理公司中某个交易台的经理，你已经注意到在纳斯达克上市的不同股票的平均买卖价差之间的差异非常巨大。

当一只股票的买卖价差和其股价的比率大于另一只股票的比率时，你的公司交易该高比率股票的成本将会变得较高。现在你做出一个假设，即纳斯达克股票百分比的买卖价差与交易该股票的做市商数量和该公司股票的市值有关。现在你想要利用多元回归分析来对该假设进行检验。

为此，你设定了一个回归模型。在该模型中因变量是百分比度量的买卖价差，而自变量分别是做市商的数量和该公司股票的市值。该回归模型使用 2002 年 12 月的 1 819 只在纳斯达克上市的股票数据进行估计。基于先前对于买卖价差所发表的研究，你将自变量和因变量都以自然对数进行表述，即采用所谓的双对数回归模型（log-log regression model）。当我们认为因变量的比例变动和自变量的比例变动间保持一个稳定的常数关系时（如我们下面所表述的），一个双对数回归模型就可能是恰当的。我们将该多元回归表述为：

$$Y_i = b_0 + b_1 X_{1i} + b_2 X_{2i} + \epsilon_i \quad (9-2)$$

式中， Y_i 表示第 i 只股票（买卖价差/股票价格）的自然对数（因变量）； X_{1i} 表示第 i 只股票纳斯达克做市商数量的自然对数； X_{2i} 表示第 i 家公司（以百万美元度量的）市场总价值的自然对数； ϵ_i 表示误差项；

在诸如式 (9-2) 的双对数回归方程中, 斜率系数被解释为弹性, 并且我们假设其为常数。例如, 在其他自变量保持不变的情况下, $b_2 = -0.75$ 意味着市值 1% 的增长, 将会造成买卖价差/股票价格 0.75% 的下降。^①

为了说明竞争程度越高造成成本越低的情况, 你需要假设做市商数量越多, 百分比买卖价差越小。因此, 你建立的第 1 个原假设和备择假设为:

$$H_0: b_1 \geq 0$$
$$H_a: b_1 < 0$$

原假设说明的是“所怀疑的”条件不成立的情况。如果证据支持拒绝原假设, 并接受备择假设, 那么你就在统计上证实了你的怀疑。^②

另外, 你认为高市值公司的股票可能具有流动性更好的市场, 因此, 其百分比买卖价差会倾向于更低一些。因此, 你建立第 2 个原假设和备择假设为:

$$H_0: b_2 \geq 0$$
$$H_a: b_2 < 0$$

对于这两个检验, 我们使用 t 检验而不是 z 检验, 因为我们不知道 b_1 和 b_2 的总体方差。^③ 假设你对于这两个检验选择了 0.01 的显著性水平。表 9-1 给出了利用 2002 年 12 月的数据所估计出的线性回归结果。

表 9-1 $\ln$ (买卖价差/股价) 和 $\ln$ (做市商数量)、 $\ln$ (市值) 的回归结果

			系数	标准误	<i>t</i> 统计量
截距			-0.758 6	0.136 9	-5.541 6
规模 (Size)			-0.279 0	0.067 3	-4.142 7
被动管理资产所占比例 (Passive management)			-0.663 5	0.024 6	-27.008 7
ANOVA	df	SS	MSS	<i>F</i>	<i>F</i> 显著性
回归	2	2 681.648 2	1 340.824 1	1 088.832 5	0.00
残差	1 816	2 236.282 0	1.231 4		
总体	1 818	4 917.930 2			
残差标准误差		1.109 7			
多元 <i>R</i> 平方		0.545 3			
观测数		1.819			

资料来源: FactSet, Nasdaq.

如果总体回归结果不显著, 那么我们将遵循不继续解释单个回归系数这个有用的原则。因此, 分析师可以先看一下 ANOVA 部分, 该部分给出了回归的总体显著性。

- ANOVA (方差分析) 部分报告了与总体解释力和回归显著性相关的数值。SS 表示平方和, MSS 表示均方和 (SS 除以 df)。 F 检验给出了回归的总体显著性。例如, F 检验一栏 0.01 的显著性说明回归在 0.01 显著性水平下显著。在表 9-1 中, 回归要比 0.01 显著性更加显著, 因为 F 检验的显著性 (在保留到小数点后两位时) 为 0。在本章的后面, 我们将介绍关于 F 检验更多的内容。

① 注意 $\Delta(\ln X) \approx \Delta X/X$, 其中 Δ 代表“变化量”而 $\Delta X/X$ 表示 X 的比例变化。我们将在例 9-11 中对于该模型进行深入讨论。
② 一个替代的有效表达为一个双边检验 $H_0: b_1 = 0$, 对应于 $H_a: b_1 \neq 0$, 该表述反映了研究者不太强的信念。具体内容参见参数检验一章。我们也可以对于我们接下去要讨论的市值进行相同的双边假设检验。
③ t 检验和 z 检验的使用在假设检验一章中已经进行了讨论。

在确定了总体回归是非常显著的之后, 分析师可以转而关注回归结果第一部分所列出的第1栏。

- 系数一栏给出了截距 b_0 和斜率系数 b_1 以及 b_2 的估计值。这些估计值都为负数, 但是它们是否显著为负呢? 标准误一栏给出了回归系数估计值的标准误 (标准差)。关于回归系数总体值假设的检验统计量具有这样的形式: (回归系数估计值 - 回归系数的假定总体值) / (回归系数的标准误)。这是一个 t 检验。在原假设下, 回归系数的假定总体值为 0。因此, (回归系数估计值 - 回归系数的假定总体值) / (回归系数的标准误) 是第3栏给出的 t 统计量。例如, 截距的 t 统计量为 $-0.7586/0.1369 = -5.5416$, 这里忽略了近似误差的影响。为了评估 t 统计量的显著性, 我们需要确定一个称为自由度 (df) 的量。^① 具体计算为: 自由度 = 观测数 - (自变量的数目) + 1 = $n - (k + 1)$ 。
- 表9-1的最后一部分给出了两个度量所估计的回归拟合程度或者解释数据好坏的指标。第1个是回归残差的标准差, 即残差标准误。这一标准差被称为估计量的标准误 (standard error of estimate, SEE)。第2个度量指标刻画的是因变量和所有自变量间联合的线性关系强弱程度。这一度量指标被称为多元 R 平方或者简单 R 平方 (因变量预测值和实际值之间相关系数的平方)。^② R^2 等于 0 意味着没有线性关系; R^2 等于 1 意味着完全线性关系。表9-1中的最后一项是样本中的观测值 (1819)。

在了解了典型回归结果中各项的含义之后, 我们可以继续来完成假设检验的讨论。

回归的估计结果支持做市商数量越多, 百分比买卖价差越小的假设: 我们拒绝 $H_0: b_1 \geq 0$, 而接受 $H_1: b_1 < 0$ 。这一结果同样也支持具有越高市值公司的股票, 其百分比买卖价差越小的假设: 我们拒绝 $H_0: b_2 \geq 0$, 而接受 $H_1: b_2 < 0$ 。为了证实在两个检验中的原假设都被拒绝, 我们可以使用书后的数值表。^③ 对于两个检验, $df = 1819 - 3 = 1816$ 。数值表没有给出那么大自由度的临界值。在 0.01 显著性水平下, $df = 200$ 的单侧检验的临界值为 2.345; 对于更大的自由度, 临界值应该比这一数值要更小。因此, 在我们的单边检验中, 如果

$$t = \frac{\hat{b}_j - b_j}{s_{\hat{b}_j}} = \frac{\hat{b}_j - 0}{s_{\hat{b}_j}} < -2.345$$

那么我们就拒绝原假设, 而接受备择假设。式中, $\hat{b}_j$ 表示回归中 b_j 的估计量, $j = 1, 2$; b_j 表示系数的假定值 (0); $s_{\hat{b}_j}$ 表示 $\hat{b}_j$ 的估计标准误。

估计量 b_1 和 b_2 的 t 值分别为 -4.1427 和 -27.0087, 都小于 -2.345。

在继续深入讨论之前, 我们应该先对于以自然对数表示的预测值进行解释。我们可以通过取反对数 (即取自然指数) 将自然对数转化为原始的单位。为了解释这一转换, 假设一只股票

① 为了计算回归中损失的自由度, 我们将自变量的数目加上 1 来对截距项加以考虑。 t 检验和自由度的概念在抽样一章中进行了讨论。

② 多元 R 平方也被称为多元决定系数, 或者简称为决定系数。

③ 参见附录 B 中 t 检验的临界值。

④ 为了有效利用表达检验统计量的记号, 在文中我们使用 b_j 来表示参数的 s 假定值 (而在其他地方, 我们用它来表示未知的总体参数)。

拥有5个纳斯达克做市商，而其公司的市值为100 000 000美元。纳斯达克做市商数量的自然对数等于 $\ln 5 = 1.6094$ ，而公司市值（以百万计）的自然对数为 $\ln 100 = 4.6052$ 。在这些值的条件下，回归模型预测买卖价差与股票价格比率的自然对数将为 $-0.7586 + (-0.2790 \cdot 1.6094) + (-0.6635 \cdot 4.6052) = -4.2632$ 。我们通过将该值作为 e 的幂次，得到其反对数的值： $e^{-4.2632} = 0.0141$ 。预测的买卖价差的值将是股票价格的1.41%。[⊖]之后，我们还将给出多元回归模型的假设；实际应用中，我们在使用一个回归估计的结果来做出预测之前，我们应该使自己确信模型的这些前提假设是得到满足的。

在表9-1中，我们给出了大部分回归软件程序所显示的常见结果。许多软件程序也会报告回归系数的 p 值。[⊗]对于每个回归系数来说， p 值是在双边检验中我们能够拒绝该系数的总体值为0的原假设的最小显著性水平。 p 值越小，拒绝原假设的证据越强。 p 值可以迅速地帮助我们确定一个自变量是否在传统的显著性水平（如0.05）或者我们认为恰当的其他标准下显著。

在估计完式(9-1)之后，我们可以写出

$$\begin{aligned}\hat{Y}_i &= \hat{b}_0 + \hat{b}_1 X_{1i} + \hat{b}_2 X_{2i} \\ &= -0.7586 - 0.2790 X_{1i} - 0.6635 X_{2i}\end{aligned}$$

式中， $\hat{Y}_i$ 表示 Y_i 的预测值，而 $\hat{b}_0$ ， $\hat{b}_1$ ， $\hat{b}_2$ 分别表示 b_0 ， b_1 ， b_2 的估计值。那么我们应该如何来解释估计出的斜率系数-0.2790和-0.6635呢？

解释多元线性回归模型中的斜率系数与相关性和回归一章中解释单个自变量回归中的斜率系数是不同的。假设我们有一个估计出的单自变量回归 $\hat{Y}_i = 0.50 + 0.72 X_{1i}$ 。对于斜率估计量0.75的解释为 X_1 每增加1单位，我们预计 Y 会相应增加0.75单位。如果我们在方程中增加第2个自变量，那么我们一般会发现 X_1 前的估计系数将不再是0.75，除非第2个自变量和 X_1 不相关。多元回归的斜率系数被称为偏回归系数（partial regression coefficients）或者偏斜率系数（partial slope coefficients），这些系数在解释时需要小心。[⊙]假设在包含第2个自变量的回归中， X_1 前的系数为0.60。我们是否可以说每1单位 X_1 的增加， Y 会相应增加0.60呢？我们没有理由这么说。对于每1单位 X_1 的增加，当 X_2 没有保持不变时，我们仍预计 Y 会增加0.75单位。而只有当第2个自变量保持不变时，我们才可以将0.60解释为对于每1单位 X_1 的增加而带来的 Y 的预计增加值。

为了解释简略的说法“第2个自变量保持不变”的含义，我们将 X_1 与 X_2 做回归，而回归的残差代表 X_1 中与 X_2 不相关的部分。于是，我们可以将 Y 与这些残差在只含一个自变量的方程中进行回归。我们会发现残差的斜率系数为0.60；通过构造，我们看到，0.60代表在移除 X_1 中与 X_2 的相关部分后每1单位 X_1 的增加所给 Y 带来的预期影响。与这一解释相同，我们也可以将0.60视为在考虑其他自变量对于 Y 预测值的影响之后，1单位 X_1 的增加给 Y 带来的预期净效应。再重申一下，偏回归系数度量的是在保持其他自变量不变的情况下，自变量1单位的增加给因变量带来的预期变动。

在表9-1的回归中应用这一解释方法，我们可以看到市值的自然对数前的估计系数为-0.6635。因此，该模型预测在保持做市商数量的自然对数不变的情况下，公司市值的自然对

⊖ 所给出的运算（取反对数） $e^{\ln X} = X$ ，使得变量恢复到以原来的单位计量的数值。

⊗ F 的显著性为0.00一栏是 F 检验的 p 值。关于 p 值更多的内容参见假设检验一章。

⊙ 这个术语来源于这样一个事实：这些系数与 Y 关于自变量的偏导数相对应。注意在这个术语中，“回归系数”指的仅仅是斜率系数。

数每增加1单位,将会带来买卖价差和股价之比的自然对数 -0.6635 的变化。我们需要注意不要认为如果两只股票所对应公司市值的自然对数间差别为1,那么这两只股票的买卖价差和股价比率的自然对数之比也会差 -0.6635 ,因为很可能两只股票的做市商数量也是不同的,而这会影响到因变量。 -0.6635 这个值是在排除了对数做市商数量对于因变量预期值的影响之后,对数市值之间差别的预期净效应。

9.2.1 多元线性回归模型的前提假设

在我们能够对多元线性回归模型做出正确的统计推断(利用普通最小二乘法对于含有多个自变量的模型进行估计)之前,我们需要了解一下关于该模型的一些假设。^①假如我们有 n 个对于因变量 Y 和自变量 $X_1, X_2, \dots, X_k$ 的观测值,而我们想要估计的等式为 $Y_i = b_0 + b_1 X_{1i} + b_2 X_{2i} + \dots + b_k X_{ki} + \epsilon_i$ 。

为了从一个多元线性回归模型中做出有效的统计推断,我们需要做出如下6个假设,以这6个假设为一组定义出了经典标准的多元线性回归模型:

1. 因变量 Y 和自变量 $X_1, X_2, \dots, X_k$ 之间满足如式(9-1)中所描述的线性关系。
2. 自变量 $(X_1, X_2, \dots, X_k)$ 不是随机的。^②同样,两个或多个自变量之间不存在完全的线性关系。^③
3. 给定自变量的条件下,误差项的期望值为0: $E(\epsilon | X_1, X_2, \dots, X_k) = 0$ 。
4. 误差项的方差对于所有观测值相同^④: $E(\epsilon_i^2) = \sigma_\epsilon^2$ 。
5. 不同观测值之间的误差项不相关: $E(\epsilon_i \epsilon_j) = 0, j \neq i$ 。
6. 误差项服从正态分布。

注意这些假设几乎与在线性回归一章中所给出的单变量线性回归模型完全相同。只不过,其中的假设2被修改为在两个或者多个自变量或者自变量的线性组合间不存在完全的线性关系。如果假设2被违背,那么我们就无法计算线性回归的估计值。^⑤同样,即使两个或者多个自变量或者自变量的线性组合间不存在完全的线性关系,线性回归也会因为两个或者多个自变量或者它们的线性组合间存在高度的相关性而遇到问题。这种高度的相关性被称为多重共线性,我们将在本章的后面讨论这个问题。我们也将讨论假定假设4和5得到满足,而事实上它们却被违背时所产生的后果。

虽然式(9-1)可能看上去只能应用于横截面数据,因为观测值的记号是相同的($i=1, \dots, n$),但是所有这些结果同样也可以应用于时间序列数据。例如,如果我们分析一家公司多个时间段的数据,我们将通常使用记号 $Y_t, X_{1t}, X_{2t}, \dots, X_{kt}$,其中第1个下标代表变量,而第2个下标代表第 t 个时间段。

① 普通最小二乘法(Ordinary least squares, OLS)是一种基于最小化回归残差平方和准则的估计方法。

② 正如相关性和回归一章中所讨论的,即使我们假设回归模型中的自变量不是随机的,但是这一假设常常并不正确。例如,标准普尔500的月度收益率就是随机的。如果自变量是随机的,那么回归模型是否就不正确了呢?幸运的是,事实并非如此。即使自变量是随机的,但其与误差项不相关,那么我们仍旧能够使用回归模型所得到的结果。例如,参见Greene(2003)或者Goldberger(1998)。

③ 没有一个自变量可以表示为另外一些自变量的线性组合。技术上来说,在这个条件中,应该包含一个等于1的常数来作为一个与截距项相对应的自变量。

④ $\text{Var}(\epsilon) = E\epsilon^2$ 和 $\text{Cov}(\epsilon_i \epsilon_j) = E(\epsilon_i \epsilon_j)$ 成立是因为 $E(\epsilon) = 0$ 。

⑤ 当我们遇到这类线性关系(称为完全共线性(perfect collinearity))时,我们无法计算线性回归估计值中所需的矩阵的逆。关于这个问题的深入描述参见Greene(2003)。

例 9-2 解释养老基金业绩表现的因素

阿姆巴奇谢尔 (Ambachtsheer)、卡佩勒 (Capelle) 和沙伊贝尔赫特 (Scheibelhut) (1998) 对于哪些因素影响养老基金的业绩表现进行了检验。具体地来说, 他们想要了解 80 家美国和加拿大养老基金的风险调整后的净附加值 (risk-adjusted net value added, RANVA) 是否依赖于每家基金的规模和基金中被动管理 (指数化) 资产的比率。利用 4 年 (1993 ~ 1996 年) 的 80 家基金的数据, 作者将 RANVA 与养老基金的规模和养老基金中被动管理的资产比率进行了回归。^① 他们使用回归式

$$\text{RANVA}_i = b_0 + b_1 \text{Size}_i + b_2 \text{Passive}_i + \epsilon_i$$

式中, RANVA_i 表示 1993 ~ 1996 年基金 i 的平均 RANVA (以百分比表示); Size_i 表示基金 i 所管理的平均资产的以 10 为底的对数值; Passive_i 表示基金 i 被动管理资产所占的比重 (以小数表示)。

表 9-2 给出了他们的分析结果。^②

表 9-2 RANVA 与规模和被动管理资产所占比例之间回归的结果

	系数	标准误	t 统计量
截距	-2.1	0.45	-4.7
规模	0.4	0.14	2.8
被动管理资产所占比例	0.8	0.42	1.9

资料来源: Ambachtsheer, Capelle, and Scheibelhut (1998)。

假如我们使用表 9-2 中的结果来对于养老基金的规模对其 RANVA 没有影响的原假设进行检验。我们的原假设为表示规模的变量前的系数等于 0 ($H_0: b_1 = 0$), 而我们的备择假设为该系数不等于 0 ($H_a: b_1 \neq 0$)。检验该假设的 t 统计量为

$$\frac{\hat{b}_1 - b_1}{s_{\hat{b}_1}} = \frac{0.4 - 0}{0.14} = 2.8$$

在 80 个观测值和 3 个系数的情况下, t 统计量具有 $80 - 3 = 77$ 个自由度。在 0.05 显著性水平下, t 的临界值约为 1.99。所计算出的规模前系数的 t 统计量为 2.8, 这表明我们有很充分的理由拒绝规模和 RANVA 不相关的原假设。估计出的系数 0.4 告诉我们, 当保持被动管理资产比例不变, 基金规模每增加 10 倍 (Size_i 每增加 1), RANVA_i 会预期相应地增加 0.4 个百分点 (40 个基点)。因为 Size_i 是以 10 为底的平均资产的对数值, Size 增加 1 单位意味着基金资产增长 10 倍。

当然, 规模和 RANVA 之间显然没有直接的因果关系: 表现更好的基金往往会吸引更多的

① 如前面的脚注所述, 技术上常把一个等于 1 的常数作为回归中截距项所对应的自变量包含在自变量之中。因为本章中给出的所有回归都包含一个截距项, 所以我们将不再在本章剩余部分对于将一个常数作为一个自变量的情况进行分开说明。

② 规模是以 10 为底的平均资产的对数值。对数变换是对于那些在很大区间内取值的自变量进行处理的常用方法; 公司规模和基金规模就是属于这种类型的两个变量。利用对数变换的一个原因是改善残差的统计性质。如果作者没有对于资产取对数, 而是直接将资产作为自变量, 那么回归模型很可能也就不能解释 RANVA。

资产。这个回归等式是与这个结果以及规模越大的基金表现越好一致的。一方面，我们可以认为规模越大的基金表现越好。另一方面，我们也可以认为表现越好的基金能够吸引更多的资产，从而使其规模变得越大。

现在假设我们想要检验被动管理资产所占比例与 RANVA 不相关的原假设；我们想要检验被动管理资产所占比例前的系数是否等于 0 ($H_0: b_2 = 0$)，其对应于被动管理资产所占比例前的系数不等于 0 的备择假设 ($H_a: b_2 \neq 0$)。检验该假设的 t 统计量为

$$\frac{\hat{b}_2 - b_2}{s_{\hat{b}_2}} = \frac{0.8 - 0}{0.42} = 1.9$$

在 0.05 显著性水平下， t 检验的临界值为 1.99，而在 0.10 显著性水平下约为 1.66。因此，在 0.10 显著性水平下，我们可以拒绝被动管理资产所占比例对于 RANVA 没有影响的原假设；然而，我们无法在 0.05 显著性水平下拒绝原假设。虽然研究者经常使用 0.05 或者更小的显著性水平，但是这些和其他类似的结果已经充分说明养老基金投资者们已经增大了对于养老基金资产被动管理的使用。我们可以将被动管理资产所占比例前的系数 0.8 解释为当保持基金规模不变，基金被动管理资产所占比例每增加 0.10，预计基金的 RANVA 会相应地增加 0.08 个百分点（8 个基点）。

例 9-3 解释富达精选科技股基金的收益率

假设你正考虑投资富达精选科技股基金（Fidelity Select Technology Fund, FSPTX），一家美国专门投资于科技股的共同基金。你想要了解该基金的表现更像大盘成长型基金还是大盘价值型基金。^①你决定估计下面的回归式

$$Y_t = b_0 + b_1 X_{1t} + b_2 X_{2t} + \epsilon_t$$

式中， Y_t 表示 FSPTX 的月度收益率； X_{1t} 表示标准普尔 500/BARRA 成长型股票指数的月度收益率； X_{2t} 表示标准普尔 500/BARRA 价值型股票指数的月度收益率。

标准普尔 500/BARRA 成长型和价值型股票指数分别代表了市场中主要的大盘成长股和价值股的表现。

表 9-3 给出了使用 1998 年 1 月到 2002 年 12 月月度数据所得到的线性回归的结果。在回归中所估计出的斜率为 0.007 9。因此，如果标准普尔 500/BARRA 成长型股票指数的收益率和标准普尔 500/BARRA 价值型股票指数的收益率在某个特定月份等于 0，那么该回归模型预测 FSPTX 的收益率将为 0.79%。大盘成长型股票指数收益率前的系数为 2.230 8，而大盘价值型股票指数收益率前的系数为 -0.414 3。因此，如果给定一个月份的标准普尔 500/BARRA 成长型股票指数收益率为 1%，而标准普尔 500/BARRA 价值型股票指数收益率为 -2%，那么该模型预测 FSPTX 的收益率将为 $0.0079 + 2.2308(0.01) - 0.4143(-0.02) = 3.85\%$ 。

① 这一回归与基于收益率的风格分析相关，这一分析是在专业投资中最为常用的回归分析应用之一。关于这方面的更多内容，参见该领域的开创者 Sharpe (1988) 以及 Buetow, Johnson 和 Runkle (2000)。

表 9-3 FSPTX 收益率和标准普尔 500/BARRA 成长型和价值型股票指数收益率间的回归结果

			系数	标准误	t 检验量
截距			0.007 9	0.009 1	0.863 5
标准普尔 500/BARRA 成长型股票指数			2.230 8	0.229 9	9.703 4
标准普尔 500/BARRA 价值型股票指数			-0.414 3	0.259 7	-1.595 3
ANOVA	df	SS	MSS	F	F 显著性
回归	2	0.864 9	0.432 4	86.448 3	5.48E-18
残差	57	0.285 1	0.005 0		
总体	59	1.150 0			
残值的标准误		0.070 7			
多元 R 平方		0.752 1			
观测数		60			

资料来源：Ibbotson Associates.

我们可能想要了解标准普尔 500/BARRA 价值型股票指数收益率前的系数是否在统计上显著。我们的原假设论述为该系数等于 0 ($H_0: b_2 = 0$)；我们的备择假设为该系数不等于 0 ($H_a: b_2 \neq 0$)。

我们对于该原假设检验所用的 t 检验构造如下：

$$t = \frac{\hat{b}_2 - b_2}{s_{\hat{b}_2}} = \frac{-0.4143 - 0}{0.2597} = -1.5953$$

式中， $\hat{b}_2$ 表示 b_2 的回归估计值； b_2 表示系数的假定值[⊖] (0)； $s_{\hat{b}_2}$ 表示 $\hat{b}_2$ 估计出的标准误。

这个回归有 60 个观测值和 3 个系数（两个自变量和一个截距项）；因此， t 检验具有 $60 - 3 = 57$ 个自由度。在 0.05 的显著性水平下，检验统计量的临界值为 2.00。[⊖] 检验统计量的绝对值为 1.5953。因为检验统计量的绝对值小于临界值 ($1.5953 < 2.00$)，所以我们无法拒绝 $b_2 = 0$ 的原假设。（注意表 9-3 中所给出的 t 检验和其他回归表中给出的一样，都是针对回归系数总体值等于 0 这个原假设进行的检验。）

类似的分析表明在 0.05 显著性水平下，我们无法拒绝截距项为 0 的原假设 ($H_0: b_0 = 0$)，所以不能接受截距项不为 0 的备择假设 ($H_a: b_0 \neq 0$)。表 9-3 给出的检验该假设的 t 统计量为 0.8635，这个结果在绝对数上小于临界值 2.00。然而，在 0.05 显著性水平下，我们可以拒绝标准普尔 500/BARRA 成长型股票指数前系数等于 0 的原假设 ($H_0: b_1 = 0$)，而接受该系数不等于 0 的备择假设 ($H_a: b_1 \neq 0$)。正如表 9-3 所示，检验该假设的 t 统计量的值为 9.70，该结果远大于临界值 2.00。因此，多元回归分析表明 FSPTX 的收益率与标准普尔 500/BARRA 成长型股票指数的收益率高度密切相关，而与标准普尔 500/BARRA 价值型股票指数的收益率不相关 (t 统计量 -1.60 在统计上不显著)。

⊖ 为了有效利用表达检验统计量的记号，在这里，我们使用 b_2 来表示参数的假定值（而在其他地方，我们用它来表示未知的总体参数）。
⊖ 参见附录 B 中 t 检验的值。

9.2.2 预测多元线性回归模型中的因变量

金融分析师们经常想要基于给定自变量的值来预测多元回归中因变量的值。我们先前已经讨论了如何对于只含有单个自变量的情况进行预测。而对于多元线性回归的预测过程与之前的讨论是非常类似的。

要利用多元线性回归模型来对于因变量的值进行预测，我们要采取这样3个步骤：

- 1. 获得回归中参数 $b_0, b_1, \dots, b_k$ 的估计量, $\hat{b}_0, \hat{b}_1 \dots \hat{b}_k$, 的值。
- 2. 确定自变量的给定值 $\hat{X}_{1i}, \hat{X}_{2i}, \dots, \hat{X}_{ki}$ 。
- 3. 利用下述等式，计算因变量 $\hat{Y}_i$ 的预测值。

$$\hat{Y}_i = \hat{b}_0 + \hat{b}_1 \hat{X}_{1i} + \hat{b}_2 \hat{X}_{2i} + \dots + \hat{b}_k \hat{X}_{ki} \tag{9-3}$$

利用估计出的回归式来预测因变量有两个实际要点需要注意。首先，我们应该确保回归模型的假设都要得到满足。其次，我们对基于模型所用来估计的数据集范围之外的自变量进行的预测要加以小心；该预测经常是不可信的。

例 9-4 预测养老基金的 RANVA

在例 9-2 中，我们基于以 10 为底的基金管理的资产规模 ($Size_i$) 和基金中被动管理资产所占比例 ($Passive_i$) 来解释美国和加拿大养老基金的 RANVA。

$$RANVA_i = b_0 + b_1 Size_i + b_2 Passive_i + \epsilon_i$$

现在我们可以利用表 9-2 中所给出的回归结果（摘录如下）来预测养老基金的业绩表现（RANVA）。

假如某个基金所管理的资产为 10 000 000 美元，而被动管理资产所占比例为 25%。以 10 为底的管理资产

表 9-2 (摘录)

	系数
截距	-2.1
规模	0.4
被动管理资产所占比例	0.8

规模等于 $\log(10\,000\,000) = 7$ 。基金被动管理资产所占比例为 0.25。相应地，基于该回归，该基金的 RANVA 预测值为 $-2.1 + (0.47) + (0.8 \times 0.25) = 0.9\%$ （90 个基点）。对于养老基金 10 000 000 美元的管理资产以及 25% 的被动管理资产比例，由该回归所预测出的 RANVA 为 90 个基点。

当利用线性回归模型对因变量进行预测时，我们遇到两类不确定性：回归模型本身的不确定性（反映在估计量的标准误中）以及回归模型参数估计值的不确定性。在相关性和回归一章中，我们给出了对单个自变量线性回归建立预测区间的具体步骤。然而，对于多元回归，计算包含这两类不确定性的预测区间需要矩阵代数，而这已经超出了本书的范围。[⊖]

9.2.3 检验是否所有回归系数为零

之前，我们介绍了如何对于单个回归系数进行假设检验。但是回归作为一个总体的显著性又是怎么样的呢？作为一个总体，自变量是否能帮助解释因变量？要回答这个问题，我们需要检验回归中的所有斜率系数同时等于 0 的原假设。在本节中，我们讨论方差分析（analysis of variance，

⊖ 更多相关内容，参见 Greene (2003)。

ANOVA), 该方法提供了关于回归解释力的信息和上述原假设的 F 检验的输入变量。

如果回归模型中的自变量没有一个能够帮助解释因变量, 那么所有的斜率系数应该都等于 0。然而, 在多元回归中, 我们无法用检验每个斜率系数等于 0 的 t 检验来对于所有斜率系数等于 0 的原假设进行检验, 因为关于单个自变量的检验无法解释自变量间的相互影响。例如, 经典的多重共线性的问题就是即使没有一个关于单个斜率系数估计量的 t 统计量是显著的, 但是我们可以拒绝所有斜率系数等于 0 的假设。相反, 我们也可以构造所估计的斜率系数显著不为 0, 而总体却不显著这样不常见的例子。

为了检验多元回归模型中所有斜率系数共同等于 0 的原假设 ($H_0: b_1 = b_2 = \cdots = b_k = 0$) 和其对应的至少有一个斜率系数不等于 0 的备择假设, 我们使用 F 检验。 F 检验可以看作对于回归总体显著性的一个检验。

要正确地计算原假设的检验统计量, 我们需要 4 个输入变量:

- 总观测数 n ;
- 要估计的回归系数的总数为 $k+1$, 其中 k 表示斜率系数的个数。
- 误差或者残差的平方和, $\sum_{i=1}^n (Y_i - \hat{Y}_i)^2 = \sum_{i=1}^n \hat{\epsilon}_i^2$, 简称为 SSE, 其也被称为剩余平方和 (不能解释的变差);[⊖]
- 回归平方和, $\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2$, 简称为 RSS。[⊖]这个数值是 Y 偏离其在回归方程所能解释均值的偏差 (能解释的变差)。

确定所有斜率系数等于 0 的 F 检验基于的是由上面列出的 4 个值所计算出的 F 统计量。[⊖] F 统计量度量的是回归方程解释因变量变动的好坏; 它是平均回归平方和和平均平方误差和之比。

我们通过将回归平方和除以所要估计的斜率系数数目 k 来计算得到平均回归平方和。我们通过将平方误差和除以观测数 n 减去 $(k+1)$ 来计算得到平均平方误差和。在这两个计算中的除数是计算 F 统计量的自由度。对于 n 个观测值和 k 个斜率系数, 所有斜率系数都等于 0 的原假设的 F 检验量记为 $F_{k, n-(k+1)}$ 。下标表明检验的分子具有 k 个自由度 (分子自由度), 而分母具有 $n-(k+1)$ 个自由度 (分母自由度)。

F 统计量的公式为

$$F = \frac{\frac{\text{RSS}}{k}}{\frac{\text{SSE}}{[n - (k + 1)]}} = \frac{\text{平均回归平方和}}{\text{平均误差平方和}} = \frac{\text{MSR}}{\text{MSE}} \quad (9-4)$$

式中, MSR 是平均回归平方和, 而 MSE 是平均误差平方和。在我们回归结果的表格中, MSR 和 MSE 是输出结果中 ANOVA 部分中的 MSS (平均平方和, mean sum of squares) 一栏下的第 1 个和第 2 个量。如果回归模型能够很好解释因变量的变动, 那么 MSR/MSE 的比率将会很大。

当回归模型中的自变量没有一个可以解释因变量的变动时, 这个 F 检验会告诉我们什么呢? 在这种情况下, 回归模型中的每个预测值 $\hat{Y}_i$ 具有因变量的均值 $\bar{Y}$, 而且回归平方和为 $\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2 = 0$ 。因此, 当自变量都无法解释因变量时, 检验原假设 (所有斜率系数都等于 0) 的 F 统计量等于 0。

⊖ 在回归结果的表格中, 这个数位于“残差”(residual)一行中的“SS”一栏。

⊖ 在回归结果的表格中, 这个数位于“回归”(regression)一行中的“SS”一栏。

⊖ 关于 F 检验更加详细的描述参见假设检验一章。

在我们计算出 F 的值之后，我们具体的统计决策方法为：如果计算出的 F 值大于给定分子和分母自由度的 F 分布的上半部分 α 的临界值，我们就在 α 显著性水平下拒绝原假设。注意我们这里使用单侧的 F 检验。[⊖]附录 D 给出了 F 检验的临界值。

我们用例 9-1 来说明这个检验，在这个例子中，我们考察纳斯达克做市商数量的自然对数和股票市值的自然对数能否解释买卖价差与股票价格比率的自然对数。假设我们令该检验的显著性水平为 $\alpha = 0.05$ （即我们在原假设为真时，错误地拒绝原假设的概率为 5%）。表 9-1（摘录如下）给出了这一回归所计算出的方差结果。

表 9-1 （摘录）

ANOVA	df	SS	MSS	F	F 显著性
回归	2	2 681. 648 2	1 340. 824 1	1 088. 832 5	0. 00
残差	1 816	2 236. 282 0	1. 213 14		
总体	1 818	4 917. 930 2			

这个模型具有两个斜率系数（ $k = 2$ ），所以在 F 检验的分子中的自由度为 2。由于样本有 1 819 个观测值，所以 F 检验分母的自由度个数为 $n - (k + 1) = 1\,819 - 3 = 1\,816$ 。残差平方和为 2 236. 282 0。回归平方和为 2 681. 648 2。因此，该模型中两个斜率系数都等于 0 的原假设的 F 检验统计量为

$$\frac{2\,681.648\,2/2}{2\,236.282\,0/1\,816} = 1\,088.832\,5$$

该检验统计量在斜率系数都等于 0 的原假设下是服从 $F_{2,1816}$ 分布的随机变量。在表中 0.05 的显著性水平下，我们可以查到第 2 栏中给出了分子为两个自由度的 F 分布的临界值。在该栏的底部，我们找到了所需的拒绝原假设的 F 检验的临界值介于 3.00 ~ 3.07 之间。[⊖]实际的 F 检验统计量为 1088.83 远大于该数值，所以我们拒绝这两个自变量的系数都为 0 的原假设。事实上，表 9-1，在“ F 显著性”一栏给出的 p 值为 0。该 p 值表明可以拒绝原假设的最小的显著性水平实际为 0。 F 统计量的值越大说明越小的错误拒绝原假设的可能性（即第一类错误）。

9.2.4 调整后的 R 平方

在相关性和回归一章中，我们给出了变差系数 R^2 ，作为回归估计对于数据拟合程度好坏的度量指标。然而，在多元线性回归中， R^2 却不太适合作为判断回归模型拟合数据好坏的度量指标（度量拟合的好坏）。回忆一下 R^2 被定义为

$$\frac{\text{总变差} - \text{不能解释的变差}}{\text{总变差}}$$

分子等于回归平方和（RSS）。因此， R^2 将 RSS 称为总平方和 $\sum_{i=1}^n (Y_i - \bar{Y})^2$ 。如果我们在模型中增加回归变量，而新的自变量解释了模型中的某些不能解释的变差，那么不能解释变差的数值会降低，而 RSS 会增加。而当新的自变量与因变量的相关性很弱，而且其不是回归中其他自变量的线性组合时，那么不能解释的变差也会降低。[⊖]结果，我们可以通过简单地往模型中增加许多

⊖ 我们使用单侧检验，因为随着回归解释力的增加，MSR 相对于 MSE 也必然会增加。
⊖ 我们看到数值所处的范围，因为分母具有超过 120 个自由度，而其小于无限。
⊖ 如果 $y = ax + bz$ （对于某个常数 a 和 b ），那么我们认为变量 y 是变量 x 和 z 的线性组合。一个变量可以是大于两个变量的线性组合。

只能解释很少的之前不能解释变差的自变量（甚至这些自变量所能解释的部分在统计上是不显著的）来使得 R^2 增加。

一些金融分析师们使用另一种度量拟合好坏的指标，称为调整后的 R 平方（adjusted R^2 ），或者 $\bar{R}^2$ 。当其他的变量增加到回归中时，该度量拟合好坏的指标不会自动增加。调整后的 R 平方通常是统计软件包产生的多元回归结果中的部分。

R^2 和 $\bar{R}^2$ 之间的关系为

$$\bar{R}^2 = 1 - \left(\frac{n-1}{n-k-1} \right) (1 - R^2)$$

式中， n 表示观测值的数目，而 k 表示自变量的数目（斜率系数的数目）。注意到如果 $k \geq 1$ ，那么 R^2 严格大于调整后的 R^2 。当一个新的自变量加入时，如果增加那个变量仅仅导致 R^2 很少的增加，那么 $\bar{R}^2$ 会降低。事实上，虽然 R^2 总是非负的，但是 $\bar{R}^2$ 可能是负的。^①如果我们使用 $\bar{R}^2$ 来比较回归模型，那么因变量必须在两个模型中以相同的方式进行定义并且用来估计模型的样本量也要是相同的，这两点非常重要。^②例如，如果因变量分别是 GDP（国内生产总值，gross domestic product）和 $\ln(\text{GDP})$ ，即使自变量是相同的，其 $\bar{R}^2$ 的值也是有差别的。我们应该进一步注意，一个高的 $\bar{R}^2$ 并不一定在包含正确变量的意义下表明回归设定得很好。^③我们这样谨慎的一个原因是，一个高的 $\bar{R}^2$ 可能只反映了用来估计回归数据集的特殊性。为了评价一个回归模型，我们需要考虑许多其他因素，正如我们在 5.1 节中所讨论的。

9.3 虚拟变量在回归中的使用

金融分析师经常需要将定性变量作为回归中的自变量。一种定性变量被称为虚拟变量（dummy variable），它在某个特定条件为真时取值为 1，而在该条件为假时取值为 0。^④例如，我们想要检验 1 月份的股票收益率是否与一年中其他剩余月份的收益率存在不同。我们就会在回归中包含一个在 1 月份取值为 1，而在一年中其他月份取值为 0 的自变量 X_{1t} 。我们所估计的回归模型为

$$Y_t = b_0 + b_1 X_{1t} + \epsilon_t$$

在这个等式中，系数 b_0 表示除了 1 月以外月份的 Y_t 的平均值，而 b_1 表示 1 月份 Y_t 的平均值和除 1 月份以外月份 Y_t 的平均值之间的差别。

我们在选择回归中虚拟变量的数目时需要小心。选取的原则是：如果我们想要区分 n 种不同的类别，那么我们需要 $n-1$ 个虚拟变量。例如，为了区分上述 1 月份和除 1 月份以外的收益率（ $n=2$ 个种类），我们就会使用一个虚拟变量（ $n-1=2-1=1$ ）。如果我们想要区分一年中 4 个季度，我们就应该包含表示一年中 4 个季度中 3 个季度的虚拟变量。如果我们错误地包含了 4 个虚拟变量而不是 3 个，那么我们会违背多元回归模型的假设 2，而且无法估计出该回归。下面的例子说明了在利用月度数据进行回归中虚拟变量的使用方法。

① 当 $\bar{R}^2$ 为负时，我们可以实际认为该值为 0。

② 参见 Gujarati (2003)。调整后的 R 平方的值依赖于样本大小。如果我们利用 R^2 来比较这两个回归模型，那么这些要点成立。

③ 参见 Mayer (1975, 1980)。

④ 不是所有的定性变量都是简单虚拟变量。例如，在三项选择模型（具有 3 种选择的模型）中，一个定性变量可能取的值为 0、1 或 2。

例 9-5 年内小公司股票收益率的月度效应

许多年来，金融分析师一直在关注股票收益率的季节性效应。[Ⓐ]特别地，分析师已经对于小公司股票收益率是否在年中各个月份不同进行了研究。假如我们想要检验一个小公司股票指数（罗素 2000 指数）的总收益率是否在各月之间存在不同。使用 1979 年 1 月（罗素 2000 数据所能获得的最早的日期）到 2002 年年底的数据来估计包含 1 个截距和 11 个虚拟变量（代表每年的前 11 个月）的回归。我们估计的等式为

$$\text{Returns}_t = b_0 + b_1 \text{Jan}_t + b_2 \text{Feb}_t + \cdots + b_{11} \text{Nov}_t + \epsilon_t$$

式中，每个月份的虚拟变量在其月份出现时取值为 1（如 $\text{Jan}_1 = \text{Jan}_{13} = 1$ ，因为第 1 个观测值是在 1 月份的），而在其他月份出现时取值为 0。表 9-4 给出了回归的结果。

表 9-4 罗素 2000 收益率关于月度虚拟变量的回归结果

		系数	标准误	t 统计量	
截距		0.030 1	0.011 6	2.590 2	
1 月		0.000 3	0.016 4	0.017 6	
2 月		-0.011 1	0.016 4	-0.675 3	
3 月		-0.021 1	0.016 4	-1.284 6	
4 月		-0.014 1	0.016 4	-0.856 8	
5 月		-0.013 7	0.016 4	-0.832 0	
6 月		-0.020 0	0.016 4	-1.216 4	
7 月		-0.040 5	0.016 4	-2.468 6	
8 月		-0.023 0	0.016 4	-1.402 5	
9 月		-0.037 5	0.016 4	-2.286 4	
10 月		-0.039 3	0.016 4	-2.396 6	
11 月		-0.005 9	0.016 4	-0.356 5	
ANOVA	df	SS	MSS	F	F 检验的显著性
回归	11	0.054 3	0.004 9	1.527 0	0.121 3
残差	276	0.892 4	0.003 2		
总数	287	0.946 7			
残差标准误		0.056 9			
多元 R 平方		0.057 4			
观测值		288			

资料来源：Ibbotson Associates.

截距 b_0 度量了 12 月股票的平均收益率，因为 12 月时回归方程中不存在虚拟变量。[Ⓑ]该回归等式估计出的 12 月的平均收益率为 3.01% ($\hat{b}_0 = 0.030 1$)。每个虚拟变量前的估计系数给出了虚拟变量所代表的那个月的收益率与 12 月收益率之间的估计偏差。因此，如果 1 月份所估计出的额外收益率比 12 月高出 0.03% ($\hat{b}_1 = 0.000 3$)。这样 1 月份收益率的预测值为 3.04% (12 月份的 3.01 + 额外的 0.03)。

Ⓐ 对于该问题的讨论，参见 Siegel (1998)。
Ⓑ 当 $\text{Jan}_t = \text{Feb}_t = \cdots = \text{Nov}_t = 0$ 时，收益率与 1 月到 11 月都不相关，所以该月份为 12 月，而回归方程也简化为 $\text{Return}_t = b_0 + \epsilon_t$ 。因为 $E(\text{Return}_t) = b_0 + E(\epsilon_t) = b_0$ ，所以截距 b_0 表示 12 月的平均收益率。

然而，这个回归的 R^2 较低，这表明年内小公司收益率的月度效应可能对于解释小公司收益率并不重要。我们使用 F 检验来解释各月虚拟变量前系数共同为 0 的原假设 ($H_0: b_1 = b_2 = \dots = b_{11} = 0$)。我们检验了小公司股票收益率各月变动间的显著性。表 9-4 给出了进行该方差检验所需要的数据。 F 检验分子的自由度为 11；分母的自由度为 $[288 - (11 + 1)] = 276$ 。回归平方和等于 0.0543，而误差平方和为 0.8924。因此，确定所有回归斜率系数是否共同等于 0 的 F 统计量为

$$\frac{0.0543/11}{0.8924/276} = 1.53$$

附录 D 给出了 F 检验的临界值。如果我们选择 0.05 的显著性水平，并查阅第 11 列（因为分子的自由度为 11），我们可以看到当分布的自由度为 120 时，该临界值为 1.87。实际分母的自由度为 276，所以 F 统计量的临界值小于 1.87（df = 120 时的临界值）但是要大于 1.79（自由度为无限时的临界值）检验统计量的值为 1.53，所以我们明显无法拒绝所有回归斜率系数共同等于 0 的原假设。

表 9-4 中给出的 F 检验量的 p 值为 0.1213，这表示我们可以拒绝原假设的最小显著性水平约为 0.12 或者 12%，这超过传统的 5% 的水平。在 11 个月份的虚拟变量中，7 月、9 月和 10 月具有绝对值超过 2 的 t 统计量。虽然这些虚拟变量前的系数在统计上是显著的，但是由于我们具有很多不显著的估计系数，所以造成我们无法拒绝每个月的收益率相等的原假设。这一检验表明回归中一些系数的显著性可能是随机变动所造成的。因此，我们可能会希望避开那些试图在不同月份给小公司股票配置不同投资权重的投资组合策略。

例 9-6 新发行的高收益债券价差的决定因素

弗里德森 (Fridson) 和格曼 (Garman) (1998) 使用 1995 ~ 1996 年的数据，检验了可能会对于解释新发行的高收益率债券和相同到期日的国库券之间初始收益率价差有帮助的变量。他们使用影响债券信用和利率风险的变量来对于收益率价差建立模型。他们的模型包含下列因素：

- 评级：穆迪高级等价类评级
- 零息债与否：虚拟变量（不是零息债 = 0，是零息债 = 1）
- BB-B 价差：收益率的差别（美林单个 B 指数减去双个 B 指数，以基点为单位）
- 高信用与否：虚拟变量（高信用 = 0，次级债 = 1）可赎回与否：虚拟变量（不可赎回 = 0，可赎回 = 1）
- 期限：到期期限（年数）
- 首次发行者：虚拟变量（不是 = 0，是 = 1）
- 承销商类型：虚拟变量（投资银行 = 0，商业银行 = 1）
- 利率变动

表 9-5 给出了作者的结果。

表 9-5 新发行高收益债券价差的多元回归模型（1995 ~ 1996 年）

	系数	标准误	t 统计量
截距	-213.67	63.03	-3.39
评级	66.19	4.13	16.02
零息债与否	136.54	32.82	4.16
BB-B 价差	95.31	24.82	3.84
高信用与否	41.46	11.95	3.47
可赎回与否	51.65	15.42	3.35
期限	-8.51	2.71	-3.14
首次发行者	25.23	10.97	2.30
承销商类型	28.13	12.67	2.22
利率变动	40.44	19.08	2.12
R 平方	0.56		
观测值	428		

资料来源：Fridson and Garman（1998）。

我们可以将弗里德森和格曼的结果汇总如下：

- 债券评级在回归模型的系数中具有最高的显著性水平。这一结果应该并不令人惊讶，因为评级指标反映了评级机构对于债券风险的估计。
- 零息债与否增加了收益率价差，因为零息债券比相同到期日的付息债券具有更高的利率风险。
- BB-B 价差会影响收益率，因为它反映了不同评级对于信用风险影响大小的不同市场评价。
- 高信用与否会影响收益率，因为次级债在违约时具有更低的偿还率。
- 可赎回与否会增加收益率，因为在收益率下降时，它限制了债券价格上涨的可能性。
- 期限事实上降低了收益率价差。可能期限具有一个负的系数，因为市场只愿意购买高质量公司的长期债券；低质量的公司必须发行较短期限的债券。
- 首次发行者必须支付一个溢价，因为市场对它们的了解不多。
- 商业银行承销的债券与投资银行承销的债券相比，具有一个额外的溢价。这很可能是因为市场认为投资银行具有一个吸引高质量公司客户的相对优势。
- 上个月国库券利率的增加会引起收益率价差的扩大，这可能是因为市场认为利率的增加将会使得发行高收益债券公司的经济前景变糟。

注意到这个回归中所有的系数都是在 0.05 显著性水平下显著的。最小的 t 统计量绝对值为 2.12。

9.4 回归假设的违背

在 2.1 节，我们给出了多元线性回归模型的假设。基于所估计的回归方程的推断必须在这些假设满足时才能成立。在对于金融数据进行回归分析时，分析师们需要有能力诊断出回归模型的假设是否被违背，并了解在假设被违背时所造成的后果，另外，他们还需要熟悉在假设被违背时所应采取的补救措施。在下面的几节中，我们将讨论 3 种回归假设被违背的情况：异方差、序列相关和多重共线性。

9.4.1 异方差

至今，我们都假定回归中不同观测值误差的方差都是相同的常数。以统计的术语来说，我们假设误差都是同方差的。然而，金融数据的误差经常是异方差的 (heteroskedastic)：不同观测值误差的方差都是不同的。在这一节中，我们将讨论异方差会对统计分析产生怎样的影响，如何来检测异方差以及如何来修正该问题。

我们可以通过两个图来了解同方差误差和异方差误差之间的区别。图 9-1 给出了因变量和自变量的值以及同方差误差模型拟合的回归直线。自变量和回归残差（标出的点和拟合回归直线之间的垂直距离）的值之间不存在系统的关系。图 9-2 给出了因变量和自变量的值以及存在异方差误差模型拟合的回归直线。在这里，可以看到一个明显的系统性关系：平均来看，回归残差随着自变量数量的增加而大幅度地增长。

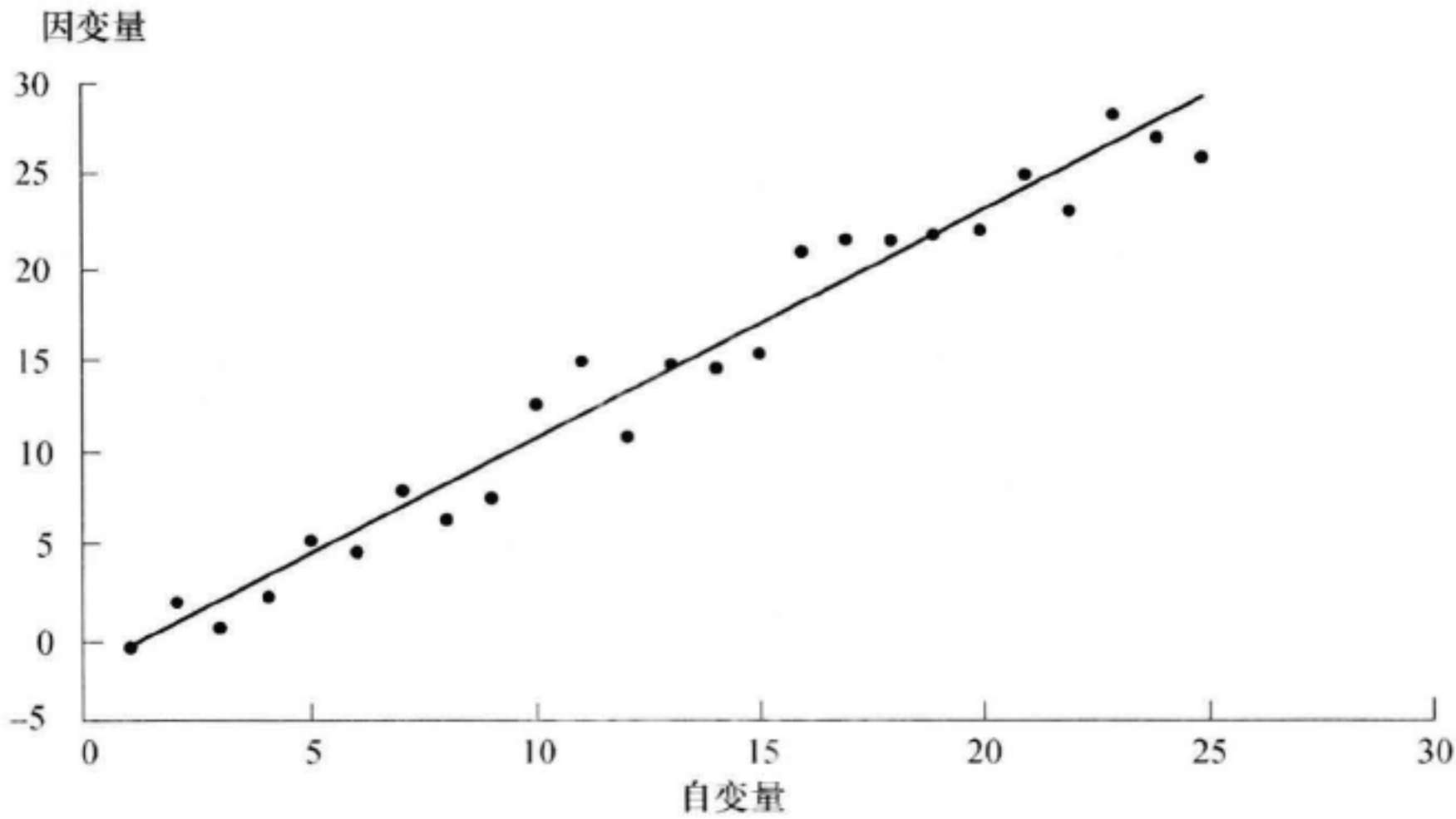


图 9-1 同方差的回归

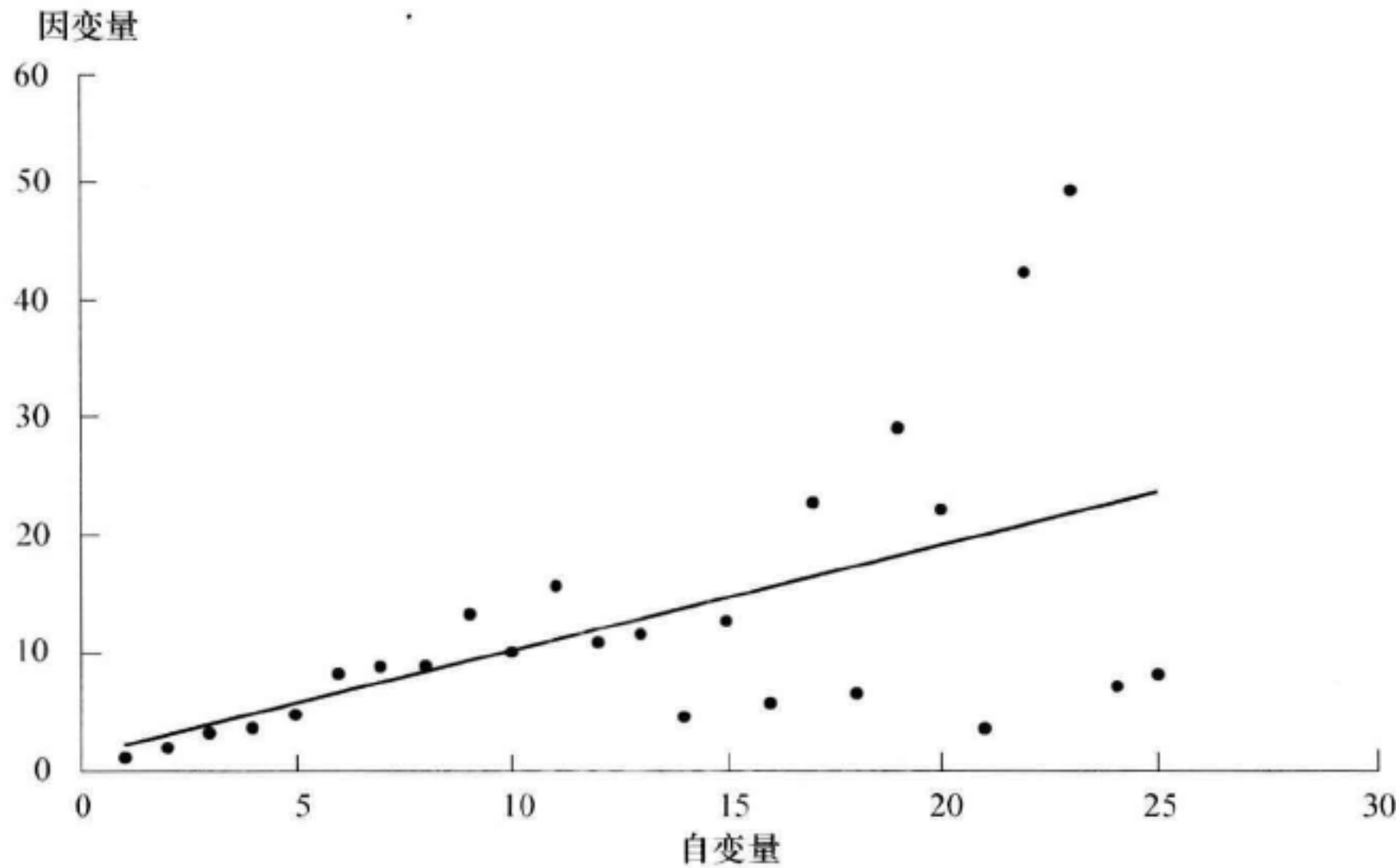


图 9-2 异方差的回归

9.4.1.1 异方差的后果

当误差同方差的假设被违背时,它所产生的后果是什么呢?虽然异方差并不影响回归参数估计量的相合性(consistency)^①,但是它会导致推断的错误。

当误差是异方差时,回归总体显著性的 F 检验是不可信的。^②而且,各个回归系数显著性的 t 检验也是不可信的,因为异方差在回归系数的标准误估计量中引入了偏差。如果一个回归表现出显著的异方差性,回归程序所计算出的标准误和检验统计量将不再正确,除非它们经过了考虑异方差的调整。

在金融数据的回归中,异方差所造成的最可能的结果是所估计出的标准误是被低估的,而 t 统计量被高估。当我们没有考虑异方差时,我们就可能发现实际上不存在的显著的关系。^③如果我们使用这些回归结果的分析来制定投资策略,所造成的实际结果可能会很严重。正如例9-7所显示的,这个问题甚至给我们对于金融模型的认识带来冲击。

例9-7 异方差和资产定价模型的检验

麦金莱(MacKinlay)和理查德森(Richardson)(1991)对于异方差在多大程度上影响了资本资产定价模型(capital asset pricing model, CAPM)的检验进行了考察。^④这些作者认为如果CAPM是正确的,那么他们应该会发现持有小公司股票和大公司股票的风险调整后收益率之间的差别是不显著的。为了实施他们的检验,麦金莱和理查德森每年将纽约和美国证券交易所的所有股票按市值分成10组,然后他们对于不同市值股票的投资组合的风险调整后收益率的系统性差别进行检验。他们估计下面的回归:

$$r_{i,t} = \alpha_i + \beta_i r_{m,t} + \epsilon_{i,t}$$

式中, $r_{i,t}$ 代表时刻 t 投资组合 i 的超额收益率(超出无风险利率的收益率) $r_{m,t}$ 代表时刻 t 整个市场全体的超额收益率

CAPM模型假设了投资组合的超额收益率可以被市场总体的超额收益率来解释。这一假设意味着对于每个 i , $\alpha_i=0$;平均来看,在考虑投资组合的系统(市场)风险之后,任何投资组合都不会获得超额收益率。

利用1926~1988年的数据以及基于等权重收益率的市场指数,当麦金莱和理查德森假设回归模型的误差服从正态分布且为同方差时,在0.05显著性水平下,他们无法拒绝CAPM。然而,他们发现当将所得到的检验统计量对于异方差进行修正之后,却能够拒绝CAPM。他们拒绝了历史数据中不存在受规模效应影响的风险调整后超额收益率的原假设。^⑤

我们已经看到异方差对于统计推断的影响可能十分严重。为了对这个概念加以更加精确的描述,我们应该区分两种宽泛的异方差——无条件的异方差和有条件的异方差。

① 通常,当回归参数估计量偏离参数真实值的概率随着回归中所使用的观测值数量的增加而减少时,我们认为回归参数的估计量是相合的,由普通最小二乘法所得到的回归参数估计量无论误差是同方差还是异方差都是相合的。具体参见抽样一章,或者参见Greene(2003)关于这个问题更加深入的讨论。

② 产生这一不可信结果的原因是平均误差平方和是给定异方差情况下真实总体方差的有偏估计量。

③ 然而,有时没有对异方差进行调整会造成过大的标准误(过小的 t 统计量)。

④ 关于CAPM更多的内容,参见例如Bodie, Kane和Marcus(2001)。

⑤ MacKinlay和Richardson也表明当使用市值加权的收益率时,无论是假设正态分布的收益率或者假设同方差还是异方差,我们都会拒绝CAPM。

当误差方差的异方差性与多元回归中的自变量不相关时,就产生了无条件异方差(unconditional heteroskedasticity)。随着这种形式的异方差违背了线性回归模型的假设4,但是它不会给统计推断产生重大的问题。

给统计推断带来主要问题的一类异方差为条件异方差(conditional heteroskedasticity),误差项方差的异方差与回归中自变量的值相关(以自变量的值为条件)。幸运的是,许多统计软件包可以方便地检测和修正条件异方差。

9.4.1.2 异方差的检验

因为条件异方差会对推断产生影响,分析师必须有能力判断出它们的存在。由于布鲁奇-帕甘(Breusch-Pagan)检验的一般性,该检验在金融研究中被广泛使用。^①

布鲁奇(Breusch)和帕甘(Pagan)(1979)提出下面对于条件异方差的检验方法:将所估计回归方程中的残差平方与回归中的自变量进行回归。如果条件异方差不存在,那么自变量将不能解释残差平方的波动。然而,如果条件异方差在原来的回归中存在的话,自变量将能解释残差平方中显著的一部分变动。自变量能解释变动是因为如果自变量能影响误差的平方,那么每个观测值的残差平方就会与自变量相关。

布鲁奇和帕甘表明在不存在条件异方差的原假设下, nR^2 (残差平方关于原始回归中自变量做回归所得到的 R^2)将是一个服从自由度等于回归中自变量数目的 χ^2 分布的随机变量。^②因此,原假设的表述为回归的误差平方项与自变量不相关。备择假设的表述为误差平方项与自变量相关。例9-8举例说明了对于条件异方差检验的布鲁奇-帕甘检验。

例9-8 利率和预期通货膨胀率间条件异方差的检验

假如一个分析师想要通过了解名义利率与预期通货膨胀率之间的关系有多密切,来确定如何在固定收益投资组合中配置资产。该分析师想要检验费雪效应(Fisher effect),该假设是由伊文·费雪(Irving Fisher)提出的,他认为预期通货膨胀率每增长1个百分点,名义利率就会增加1个百分点。^③费雪效应假设了下面的名义利率,实际利率和预期通货膨胀率之间的关系:

$$i = r + \pi^e$$

式中, i 代表名义利率; r 代表实际利率; π^e 代表预期通货膨胀率。

为了利用时间序列数据来对于费雪效用进行检验,我们可以对于名义利率设定下面的回归模型:

$$i_t = b_0 + b_1 \pi_t^e + \epsilon_t \quad (9-5)$$

注意到费雪效应预期通货膨胀率变量前的系数为1,所以我们可以陈述原假设和备择假设如下:

$$H_0: b_1 = 1$$

$$H_a: b_1 \neq 1$$

① 一些其他的检验需要更多的关于异方差函数形式的具体假设。更多内容参见 Greene (2003)。

② 布鲁奇-帕甘检验在大样本时服从一个 χ^2 分布的随机变量。这里,在技术上与回归中的截距项有关的常数1在计算自变量个数时不被计算入内。更多关于布鲁奇-帕甘检验的内容,参见 Greene (2003)。

③ 关于更多费雪效应的内容,参见 Mankiw (2000)。

我们可能会对该检验设定 0.05 的显著性水平。在我们估计式 (9-5) 之前,我们必须确定如何来度量预期通货膨胀率 (π_t^e) 和名义利率 (i_t)。

职业预测者的调查 (the survey of profession forecasters, SPF) 收集了每个季度职业预测家对于通货膨胀率的预测数据。^①我们使用这些数据作为我们度量预期通货膨胀率的指标。我们使用 3 个月的国库券收益率作为我们对于无风险名义利率的度量指标。^②我们使用 1968 年第 4 季度至 2002 年第 4 季度的季度数据来估计式 (9-5)。表 9-6 给出了回归的结果。

表 9-6 国库券收益率对于预测通货膨胀的回归结果

	系数	标准误	t 检验量
截距	0.030 4	0.004 0	7.688 7
预期通货膨胀率	0.877 4	0.081 2	10.809 6
残值的标准误		0.0220	
多元 R 平方		0.4640	
观测数		137	
杜宾-沃特森统计量		0.4673	

资料来源: Federal Reserve Bank of Philadelphia, U. S. Department of Commerce.

为了对数据是否支持费雪效应做出一个统计结论,我们计算下面的 t 统计量,然后我们将其与临界值进行比较。

$$t = \frac{\hat{b}_1 - b_1}{s_{\hat{b}_1}} = \frac{0.877\,4 - 1}{0.081\,2} = -1.509\,9$$

t 统计量为 -1.509 9,而自由度为 $137 - 2 = 135$ 的临界 t 值约为 1.98。如果我们进行的是有效的检验,那么我们在 0.05 显著性水平下无法拒绝这个回归中的真实系数为 1 且费雪效应成立的原假设。 t 检验假设误差是同方差的。因此,在我们接受 t 检验的有效性之前,我们应该检验一下误差是否是条件异方差的。如果这些误差被证明是条件异方差的,那么该检验是无效的。

我们对于费雪效应回归中的残差平方执行检验条件异方差的布鲁奇-帕甘检验 (Breusch-Pagan test)。该检验将残差平方与预期通货膨胀率进行回归。残差平方回归的 R^2 为 0.1651。由该回归所得到的检验统计量 nR^2 为 $1370.1651 = 22.619$ 。在不存在条件异方差的原假设下,该检验统计量是一个自由度为 1 的 χ^2 随机变量 (因为只有一个自变量)。

我们只有对较大的检验统计量的值才应该考虑异方差。因此,我们应该使用单侧检验来确定我们是否应该拒绝原假设。附录 C 给出了在 0.05 显著性水平下服从自由度为 1 的 χ^2 分布的检验统计量的临界值为 3.84。由布鲁奇-帕甘检验得到的检验统计量为 22.619,所以我们能在 0.05 显著性水平下拒绝不存在条件异方差的原假设。事实上,我们甚至可以在 0.01 显著性水平下拒绝不存在条件异方差的原假设,因为这种情况下检验统计量的临界值为 6.63。因此,我们认为费雪效应回归中的误差项是存在条件异方差的。原始回归中所计算得到的标准误是不正确的,因为它没有考虑异方差。因此,我们不能认为该 t 检验是有效的。

① 对于这个例子,我们使用年化的当季度 GDP 平减指数 (在 1992 年之前使用 GNP 平减指数) 增长率的 SPF 预测的中位数。
② 我们关于国库券收益率的数据是基于二级市场 3 个月国库券的收益率而得到的。因为这些收益率是以折价基准报价的,我们将它们转化为复合的年利率,因此它们是以和通货膨胀率的预测数据相同的基准进行度量的。这些收益率是无风险的,因为它们在季初就已经为人们所知,而且它们不存在违约风险。

在例 9-8 中，我们得到的结论是我们用来检验费雪效应的 t 检验不是有效的。那是否意味着我们不能使用一个回归模型来检验费雪效应呢？幸运的是，事实不是这样的。我们有方法来对回归系数的标准误所存在的异方差进行修正。使用 $\hat{b}_1$ 调整的标准误，我们可以重新进行 t 检验。正如我们将在下一节中会看到的，使用这个有效的 t 检验，我们仍旧无法拒绝例 9-8 中的原假设。

9.4.1.3 异方差的修正

金融分析师需要了解如何修正异方差的问题，因为这一修正可能会改变对于某一假设检验的结论，并最终影响到某个投资决策。（例如，在例 9-7 中，麦金莱和理查德森在将他们模型的显著性检验对于异方差进行修正之后，改变了他们的投资结论。）

我们能使用两种不同方法来修正线性回归模型中条件异方差的效应。第 1 种方法，通过计算稳健的标准误（robust standard errors）来修正存在条件异方差情况下的线性回归模型中估计系数的标准误。第 2 种方法，通过广义最小二乘法（Generalized least squares）来修改原始回归方程以消除异方差的影响。新的、修改好的模型是在不存在异方差问题的假设下进行估计的。^①这两种修正条件异方差方法背后的技术细节已经超出了本书的范围。^②然而，许多统计软件包能够方便地计算出稳健的标准误，我们推荐使用这些软件。^③

回到例 9-8 中关于费雪效应的问题，回忆一下，我们的结论是误差的方差是存在异方差的。如果我们对于回归系数标准误的条件异方差进行了修正，那么我们就得到了表 9-7 所给出的结果。在将表 9-7 和表 9-6 中的标准误进行比较时，我们会发现截距的标准误变化很小，但是预期通货膨胀率前系数（斜率系数）的标准误却增长了约 25.5%（从 0.0812 增长到 0.1019）。我们同样也注意到两张表中的回归系数是相同的，因为表 9-7 中的结果仅仅修正了表 9-6 中的标准误。

表 9-7 国库券收益率关于预期通货膨胀率回归的结果
(标准误对于条件异方差进行了修正)

	系数	标准误	t 统计量
截距	0.0304	0.0038	8.0740
预期通货膨胀率	0.8774	0.1019	8.6083
残差标准误		0.0220	
多元 R 平方		0.4640	
观测数		137	

使用 $\hat{b}_1$ 稳健的标准误，我们现在就能对于斜率系数的真实值为 1 的原假设进行有效的 t 检验了。我们发现 $t = (0.8774 - 1) / 0.1019 = -1.2031$ 。在绝对值上，该数值仍远小于拒绝斜率等于 1 的原假设所需的临界值 1.98。^④因此，在这个特殊的例子中，即使条件异方差在统计上是显著的，将其修正后仍不会使我们对于预期通货膨胀率系数的假设检验结果产生任何影响。然而，在其他的情况下，我们的统计决策可能会由于在 t 检验中使用了稳健的标准误而发生改变。关于 CAPM 的检验例 9-7 就是这样一个例子。

9.4.2 序列相关

一个比同方差假设被违背的更加常见的（可能造成更严重影响的）问题为不同观测值间回归误差不相关的假设被违背。试图解释几个时间段中的某个金融关系是有风险的，因为金融回归模型中的误差经常在时间上是相关的。

① 广义最小二乘法需要经济计量学的专门技术来对于金融数据进行准确地处理。参见 Greene (2003), Hansen (1982) 以及 Keane 和 Runkle (1998)。
② 关于这两种方法的更多内容，参见 Greene (2003)。
③ 稳健的标准误也被称为异方差一致标准误（heteroskedasticity-consistent standard errors）或者怀特修正的标准误（White-corrected standard errors）。
④ 记住，这是一个双侧检验。

当不同观测值的回归误差是相关的时候,我们认为它们是序列相关的 (serially correlated) (或者自相关)。序列相关主要出现在时间序列回归之中。在本节中,我们讨论序列相关的 3 个方面:对于统计推断的影响,如何检验以及修正的方法。

9.4.2.1 序列相关的后果

与异方差一样,线性回归中序列相关所造成的主要问题是由统计软件包所估计的回归系数标准误是不准确的。只要不存在自变量是因变量的滞后项 (前一期因变量的值),那么估计出的参数本身将是相合的,也不需要对于序列相关的影响进行调整。然而,如果某个自变量是因变量的滞后值,例如,如果前一期的国库券收益率是费雪效应回归中的一个自变量,那么误差项中的序列相关将会造成线性回归中所有参数估计值不具有相合性,而且这些估计值也不再是真实参数的有效估计量。^①

在本章中,我们假设自变量都不是因变量的滞后值。当这一假设不成立时,序列相关就会对回归系数的标准误造成影响。我们将在正序列相关的情况下对该问题进行检验,因为这种情况是最常见的。正序列相关 (positive serial correlation) 是指在序列中一个观测值正误差会增加另一个观测值正误差的可能性。正序列相关也意味着一个观测值的负的误差会增加另一个观测值负的误差的可能性。^②在检验正序列相关时,一般我们所作出的假设为对序列相关采用一阶序列相关 (first-order serial correlation),或者相邻观测值间的序列相关的形式。在时间序列中,该假设表示误差项的符号在一期和另一期之间保持不变。

虽然正序列相关不影响估计出的回归系数的相合性,但是它会影响我们进行有效统计检验的能力。首先,用来检验回归总体显著性的 F 统计量可能会被高估,因为平均误差平方和 (MSE) 将倾向于低估总体误差的方差。其次,正序列相关通常会造普通最小二乘法 (OLS) 所获得的回归系数的标准误低估真实的标准误。结果,如果回归中表现出正序列相关,那么标准线性回归分析将误导我们计算出虚假的较小的回归系数的标准误。这些较小的标准误将造成估计出的 t 统计量被高估,从而表现出本来可能不存在的显著性。这个被高估的 t 统计量可能最终导致我们比原来使用正确估计出的标准误,更大可能性地错误地拒绝关于回归模型中参数总体值的原假设。这种第一类错误会导致我们得出不恰当的投资建议。^③

9.4.2.2 序列相关的检验

我们可以选择多种检验回归模型中序列相关的方法^④,但是最常用的方法是采用由杜宾 (Durbin) 和沃特森 (Watson) (1951) 所提出的统计量;事实上,许多统计软件包会自动计算杜宾 - 沃特森 (Durbin-Watson) 统计量。计算杜宾 - 沃特森检验统计量的等式为

$$DW = \frac{\sum_{t=2}^T (\hat{\epsilon}_t - \hat{\epsilon}_{t-1})^2}{\sum_{t=1}^T \hat{\epsilon}_t^2} \quad (9-6)$$

式中, $\hat{\epsilon}_t$ 表示时刻 t 的回归残差。

① 我们将在时间序列分析一章中讨论这个问题。

② 相反,在序列负相关的情况下,一个观测值正的误差会增加另一个观测值负的误差的可能性,而一个观测值负的误差会增加另一个观测值正的误差的可能性。

③ 如果回归中存在序列负相关,那么 OLS 标准误不一定低估实际的标准误。

④ 参见 Greene (2003) 关于序列相关检验更深入的探讨。

我们可以将该等式重写为

$$\frac{\frac{1}{T-1} \sum_{t=2}^T (\hat{\epsilon}_t^2 - 2 \hat{\epsilon}_t \hat{\epsilon}_{t-1} + \hat{\epsilon}_{t-1}^2)}{\frac{1}{T-1} \sum_{t=1}^T \hat{\epsilon}_t^2} \approx \frac{\text{Var}(\hat{\epsilon}_t) - 2\text{Cov}(\hat{\epsilon}_t, \hat{\epsilon}_{t-1}) + \text{Var}(\hat{\epsilon}_{t-1})}{\text{Var}(\hat{\epsilon}_t)}$$

如果误差的方差随时间是不变的, 那么我们预期 $\text{Var}(\hat{\epsilon}_t) = \hat{\sigma}_\epsilon^2$ (对于所有的 t 成立), 其中我们使用 $\hat{\sigma}_\epsilon^2$ 来代表误差方差的常数估计值。另外, 如果误差也不是序列相关的, 那么我们预期 $\text{Cov}(\hat{\epsilon}_t, \hat{\epsilon}_{t-1}) = 0$ 。在这种情况下, 杜宾-沃特森统计量大约等于

$$\frac{\hat{\sigma}_\epsilon^2 - 0 + \hat{\sigma}_\epsilon^2}{\hat{\sigma}_\epsilon^2} = 2$$

这个等式告诉我们, 如果误差是同方差且不是序列相关的, 那么杜宾-沃特森统计量应该接近于 2。因此, 我们可以通过检验杜宾-沃特森统计量是否显著不同于 2 来对于误差不是序列相关的原假设进行检验。

如果样本很大, 那么杜宾-沃特森统计量将近似等于 $2(1-r)$, 其中 r 表示本期和上一期回归残差的样本相关系数。该估计是非常有用的, 因为它表明对于不同程度的序列相关水平, 杜宾-沃特森统计量的值是不同的。杜宾-沃特森统计量可以取 0 (在序列相关系数等于 +1 的情况下) 到 4 (在序列相关系数等于 -1 的情况下) 之间的值:

- 如果回归中不存在序列相关, 那么回归残差在时间上不相关, 而杜宾-沃特森统计量的值将等于 $2(1-r) = 2$ 。
- 如果回归残差是序列正相关的, 那么杜宾-沃特森统计量将小于 2。例如, 如果误差的序列相关系数为 1, 那么杜宾-沃特森统计量将为 0。
- 如果回归残差是序列负相关的, 那么杜宾-沃特森统计量将大于 2。例如, 如果误差的序列相关系数为 -1, 那么杜宾-沃特森统计量将为 4。

回到检验费雪效应的例 9-8 中, 表 9-6 中对于 OLS 回归的 D 杜宾-沃特森统计量为 0.4673。这一结果表明回归残差是序列正相关的:

$$\begin{aligned} DW &= 0.4673 \\ &\approx 2(1-r) \\ r &\approx 1 - DW/2 \\ &= 1 - 0.4673/2 \\ &= 0.766 \end{aligned}$$

由于序列正相关, 这一结果引起我们对于 OLS 标准误可能是不正确的关注。那么观测到的杜宾-沃特森统计量 (0.4673) 是否提供了充分地拒绝不存在序列正相关的原假设的证据呢?

如果杜宾-沃特森统计量低于一个临界值 d^* , 那么我们应该拒绝不存在序列正相关的原假设。然而, 不幸的是杜宾和沃特森告诉我们在给定样本的情况下, 我们无法知道确切的临界值 d^* 。事实上, 我们只能确定 d^* 或者位于两个值 d_u (上界值) 和 d_l (下界值) 之间, 或者位于这两个值之外。^①图 9-3 根据杜宾-沃特森统计量的相关结果, 画出了 d^* 的上界和下界值。

① 附录 E 列表给出了 0.05 显著性水平下不同估计参数数目 ($k=1, 2, \dots, 5$) 和时间区间数目 (在 15 到 100 之间) 的 d_l 和 d_u 的值。



图 9-3 杜宾 - 沃特森统计量的值

从图 9-3 中，我们得到如下结果：

- 当杜宾 - 沃特森（DW）统计量小于 d_l 时，我们拒绝序列不正相关的原假设。
- 当杜宾 - 沃特森（DW）统计量落入 d_l 和 d_u 之间时，该检验无法得出确定性的结果。
- 当杜宾 - 沃特森（DW）统计量大于 d_u 时，我们无法拒绝序列不正相关的原假设。[⊖]

回到例 9-8 中，费雪效应的回归具有一个自变量和 137 个观测值。杜宾 - 沃特森统计量为 0.4673。如果我们查阅附录 E 中 $k = 1$ 的一栏，我们会看到在 0.05 显著性水平下应该拒绝不相关的原假设，而接受序列正相关的备择假设，因为杜宾 - 沃特森统计量远低于 $k = 1$ 和 $n = 100$ 的临界值 d_l (1.65)。 d_l 的水平在 137 个观测值的样本中会更高。这一显著序列正相关的发现告诉我们，回归中的 OLS 标准误很可能显著低估了真实的标准误。

9.4.2.3 序列相关的修正

当一个回归具有显著的序列相关时，我们有两种补救的措施。第 1 种，我们将系数的标准误调整为考虑序列相关的线性回归参数的估计值。第 2 种，我们对于回归方程本身进行修改以消除序列相关。我们推荐使用第 1 种方法来处理序列相关的问题；第 2 种方法可能会导致不相合的参数估计值，除非处理得非常小心。

最为人们所普遍使用的调整标准误的方法是汉森（Hansen）（1982）所提出的，该方法在许多统计软件包中都是一个标准的结果。[⊖]汉森方法的一个额外优点是它同时修正了条件异方差。[⊗]表 9-8 给出了利用汉森方法将表 9-6 中的标准误对于序列相关和异方差进行了调整的结果。注意到截距和斜率系数都与之前的回归中完全相同。然而，稳健的标准误现在更大了，几乎是 OLS 标准误的两倍。因为回归误差中严重的序列相关性，OLS 大大地低估了回归中估计参数的不确定性。

我们同样注意到杜宾 - 沃特森统计量和表 9-6 相比没有发生变化。虽然序列相关还没有消除，但是标准误已经对于序列相关进行了调整。

现在假设我们想要检验我们最初的预期通货膨胀率前的系数等于 1 的原假设（费雪效应）（ $H_0: b_1 = 1$ ）对应于通货膨胀率前的系数不等于 1 的备择假设（ $H_a: b_1 \neq 1$ ）。使用修正的标准误，关于该原假设的检验统计量的值为

表 9-8 国库券收益率关于预期通货膨胀率进行回归的结果
(标准误对于条件异方差和序列相关已经进行了调整)

	系数	标准误	t 统计量
截距	0.0304	0.0069	4.4106
预期通货膨胀率	0.8774	0.1729	5.073 0
残差标准误		0.0220	
多元 R 平方		0.4640	
观测数		137	
杜宾 - 沃特森统计量		0.4673	

资料来源：Federal Reserve Bank of Philadelphia, U. S. Department of Commerce.

⊖ 当然，有时回归模型中的序列相关系数为负而不是正。对于一个序列不相关的原假设，如果 $DW < d_l$ （表明显著的序列正相关），或者如果 $DW > 4 - d_l$ （表明显著的序列负相关），原假设都会被拒绝。

⊗ 这一修正方法有许多不同的名称，包括一致序列相关标准误——序列相关和异方差调整后的标准误，稳健的标准误，Hansen-White 标准误。分析师可能也会说他们使用 Newey-West 方法来计算稳健的标准误。

⊘ 我们不经常使用 Hansen 方法来修正序列相关和异方差，这是因为有时回归的误差不是序列相关的。

$$\frac{\hat{b}_1 - b_1}{s_{\hat{b}_1}} = \frac{0.8774 - 1}{0.1729} = -0.7091$$

0.05 和 0.01 显著性水平下的临界值都远大于 0.7091 (t 值的绝对值), 因此我们无法拒绝原假设。

在这个例子中, 我们对于费雪效应的结论并不受到序列相关的影响, 但是在考虑序列相关和条件异方差之后的斜率系数的标准误 (0.1729) 大于 OLS 标准误 (0.0812) 两倍多。因此, 对于一些假设, 序列相关和条件异方差可能对于我们是否接受或者拒绝原假设产生重大的影响。[⊖]

9.4.3 多重共线性

多元线性回归模型的第2个假设是两个或者更多自变量之间不存在完全的线性关系。当其中的一个自变量完全是其他自变量的线性组合时, 软件就无法对于该回归进行估计了。这种情况被称为完全共线性, 与多重共线性相比, 它在实际中较少受人们关注。[⊖]当两个或者更多自变量 (或者自变量的线性组合) 之间高度相关 (但不是完全相关), 多重共线性 (multicollinearity) 问题就产生了。在多重共线性的情况下, 我们可以对回归进行估计, 但是回归结果的解释会存在问题。多重共线性在实践中受到高度的关注, 因为金融变量间的近似线性关系是非常普遍的。

9.4.3.1 多重共线性的后果

虽然多重共线性的问题并不影响回归系数的 OLS 估计量的相合性, 但是估计量会变得极端不准确而且不可信。另外, 它使得我们无法区分单个自变量对于因变量的影响。这些后果都反映在高估的回归系数的 OLS 标准误中。由于高度的标准误, 系数的 t 检验具有很小的势 (拒绝原假设的能力)。

9.4.3.2 多重共线性的侦测

不同于异方差和序列相关的例子, 我们将不提供正式的检验多重共线性的统计检验。在实务中, 多重共线性经常只是存在程度大小的问题, 而不是存在与不存在的问题。[⊖]

分析师应该注意, 虽然有时我们会建议使用自变量两两之间的相关程度来评估多重共线性的程度, 但是一般这是不充分的。虽然自变量两两间很高的相关系数可以反映出多重共线性, 但是在多重共线性问题存在时, 自变量两两间的相关系数并不一定很高。[⊖]换句话说, 自变量两两间较高的相关系数并不是多重共线性的必要条件, 而自变量两两间较低的相关系数并不意味着不存在多重共线性的问题。自变量相关系数可以作为回归中是否存在多重共线性的情况的一个合理判断的情况是只有在两个自变量的回归中才会出现。

多重共线性的经典症状是即使估计出的斜率系数的 t 统计量是不显著的, 回归也会具有很高 R^2 (显著的 F 统计量)。不显著的 t 统计量反映出高估的标准误。虽然系数可能估计得非常不准

⊖ 序列相关也会影响预测的准确性。我们将在时间序列一章中讨论这个问题。

⊖ 为了给出一个完全共线性的例子, 假设我们试图用一个包含净销售额、销售成本和销售利润作为自变量的回归来解释公司信用评级。因为根据定义, 销售利润 = 净销售额 - 销售成本, 所以这些变量间存在完全线性关系。这类错误是相对明显的 (也容易避免)。

⊖ 参见 Kmenta (1986)。

⊖ 即使自变量之间的相关系数很小, 也会存在自变量的线性组合间较高的相关性, 从而产生多重共线性的问题。

确（如较低的 t 统计量所反映的），但是自变量作为一个整体却可以较好地解释因变量，而很高的 R^2 反映了这一情况。例 9-9 给出了这一诊断。

例 9-9 在解释富达精选科技股基金收益率中所存在的多重共线性

在例 9-3 中，我们将富达精选科技股基金（FSPTX）的收益率与标准普尔 500/BARRA 成长型股票指数的收益率和标准普尔 500/BARRA 价值型股票指数的收益率进行了回归。表 9-9 给出了我们回归的结果，其使用了 1998 年 1 月到 2002 年 12 月的数据。成长型股票指数收益率的 t 统计量为 9.703 4，远大于 2.0，这表明成长型股票指数前的系数在标准显著性水平下显著不为 0。另一方面，价值型股票指数收益率的 t 统计量为 -1.595 3，因此它在统计上不显著。这一结果表明 FSPTX 的收益率和成长型股票指数的收益率相关而与价值型股票指数的收益率关系不密切。然而，成长型股票指数前的系数为 2.23。这一结果表示 FSPTX 收益率的波动性远大于成长型股票指数。

表 9-9 FSPTX 收益率和标准普尔 500/BARRA 成长型股票指数收益率和价值型股票指数收益率进行了回归的结果

			系数	标准误	t 统计量
截距			0.007 9	0.009 1	0.863 5
标准普尔 500/BARRA 成长型股票指数			2.230 8	0.229 9	9.703 4
标准普尔 500/BARRA 价值型股票指数			-0.414 3	0.259 7	-1.595 3
ANOVA	df	SS	MSS	F	F 检验的显著性
回归	2	0.8649	0.4324	86.4483	5.48E-18
残差	57	0.2851	0.0050		
总数	59	1.1500			
残差标准误		0.0707			
多元 R 平方		0.7521			
观测值		60			

资料来源：Ibbotson Associates.

同样注意到这一回归显著地解释了大部分 FSPTX 收益率的变动量。具体来说，这个回归的 R^2 为 0.7521。因此，几乎 75% 的 FSPTX 收益率的变动被标准普尔 500/BARRA 成长型股票指数和价值型股票指数的收益率所解释。

现在假设我们将标准普尔 500 指数本身的收益率加入到标准普尔 500/BARRA 成长型股票指数和价值型股票指数的收益率之中，并再次进行一个线性回归。标准普尔 500 包含了这两种风格指数的成分股，因此我们引入了一个严重的多重共线性的问题。

表 9-10 给出了这个回归的结果。注意到这个回归中 R^2 相对于之前回归中 R^2 间的改变几乎难以察觉（从 0.752 1 增长到 0.753 9），但是现在系数的标准误更大了。在前一个回归中加入标准普尔 500 收益率并不能比前一个回归解释任何更多的 FSPTX 收益率的变动，但是现在所有的系数都在统计上不显著了。这是在文中提到的典型的多重共线性的情况。

表 9-10 FSPTX 收益率关于标准普尔 500/BARRA 成长型股票指数和价值型股票指数收益率以及标准普尔 500 指数收益率的回归结果

			系数	标准误	t 统计量
截距			0.007 2	0.009 2	0.776 1
标准普尔 500/BARRA 成长型股票指数			-1.132 4	5.244 3	-0.215 9
标准普尔 500/BARRA 价值型股票指数			-3.491 2	4.800 4	-0.727 3
标准普尔 500 指数			6.443 6	10.0380	0.641 9
ANOVA	df	SS	MSS	F	F 检验的显著性
回归	3	0.867 0	0.289 0	57.175 1	4.73E-17
残差	56	0.283 0	0.005 1		
总数	59	1.150 0			
残差标准误			0.071 1		
多元 R 平方			0.753 9		
观测值			60		

资料来源：Ibbotson Associates.

甚至当我们无法观测到经典的 t 统计量不显著但是 F 检验高度显著的症状时，多重共线性仍可能是一个潜在的问题。高级的教科书会提供更加深入的工具来帮助我们诊断多重共线性。[⊖]

9.4.3.3 多重共线性的修正

最直接的修正多重共线性的方法是排除一个或者多个回归变量。在上面的例子中，我们可以看到如果标准普尔 500/BARRA 成长型股票指数和价值型股票指数的收益率被包含在回归中时，标准普尔 500 全收益率就不应该包含在回归中，因为全部标准普尔 500 指数的收益率是成长型股票收益率和价值型股票收益率的加权平均值。然而，在许多情况下，要解决多重共线性的问题并不简单，而你将需要通过加入或者排除不同自变量的多次尝试来确定多重共线性问题的来源。

9.4.4 异方差、序列相关、多重共线性：问题的总结

我们已经讨论了异方差、序列相关和多重共线性，可能会在解释回归结果时所产生的问题。我们所提到的这些对于回归假设的违背都会使我们在做出有效推断的出现问题。所以分析师应该在解释统计检验结果之前确认模型的假设是否都得到满足。

表 9-11 给出了这些问题、它们对线性回归模型所产生的影响（分析师可以使用回归软件来看到这些影响）以及解决这些问题方法的一个汇总。

表 9-11 线性回归中的问题及其解决方法

问题	影响	解决方法
异方差	不正确的标准误	使用稳健的标准误（对条件异方差进行修正的）
序列相关	不正确的标准误（如果因变量的滞后值被用作自变量会有其他问题产生）	使用稳健的标准误（对序列相关进行修正的）
多重共线性	高 R^2 和低 t 统计量	去除一个或者多个自变量；通常没有理论的解决方法。

⊖ 参见 Greene (2003)

9.5 模型设定和设定中的错误

至今,我们都假设我们估计的回归模型无论怎样都是正确设定的。模型设定 (Model specification) 指的是对于回归中所包含的变量和回归方程函数形式的设定。下面,我们首先给出一些宽泛的正确设定回归的准则,然后我们转向3种错误的模型设定、函数形式的错误设定、回归量和误差项相关以及其他一些有关时间序列的设定错误。这些设定错误中的每一种都会使得利用OLS所得到的统计推断归于无效;这些错误设定中的大部分都会造成所估计出的回归系数不是相合的。

9.5.1 模型设定的原则

在讨论模型设定的原则时,我们需要知道存在着相互冲突的关于如何进行模型设定的观点。而且,我们使用回归分析的目的可能影响我们对于模型设定的选择。然而,下面给出的原则一般得到人们相对广泛的应用。

- 模型应该扎根于具有说服力的经济理论。我们应该能够提供变量选择背后的经济学理由,并且该理由应该有充分的道理。当这个条件满足时,我们增加了模型在新数据条件下产生准确预测值的可能性。这一方法和在抽样一章变量选择过程中的数据挖掘不同。在数据挖掘中,研究者本质上是通过最大化利用特定数据集的特征来构建模型。
- 给定变量的性质,为模型中的变量所选择的函数形式应该恰当。作为一个例子,我们只通过基金收益率和市场收益率来研究共同基金的市场择时问题。有人可能会认为对于成功的择时者,共同基金的收益率相对于市场收益率应该是弯曲的,因为当市场收益率较高(低)时,一个成功的择时者将倾向于增加(减少)其贝塔。模型设定应该反映预期的非线性关系。[⊖]在其他情况中,我们可能将数据进行转换以使得回归模型的假设得以较好的满足。
- 模型应该要简洁。在这里,“简洁”表示要以最少的设定来达到尽量多的目的。我们应该希望回归中所包含的每个变量都能发挥各自不可缺少的作用。
- 在接受模型前,应该对于其是否违背回归假设进行检验。我们已经讨论了检测异方差、序列相关和多重共线性是否存在的方法。在这类诊断结果的基础上,我们可以得出是否需要修改所包含的变量和(或)它们的函数形式的设定。
- 在接受模型前,我们要对模型进行检验,并判断模型在样本外是否有用。术语“样本外”是指在回归估计所用的数据集之外的观测值。一个貌似真实的模型可能在样本外表现不佳,因为经济关系在样本区间之外发生了改变。这个可能性本身是需要我们去了解的。然而,第2种解释可能是经济关系没有发生改变,而只是因为模型只能解释这一特定的数据集。

在给出了这些宽泛的关于模型设定的原则之后,我们回到前面对于特定模型设定误差的讨论。了解这些误差会帮助分析师建立更好的模型,同时使得投资研究的使用者们获得更多的

⊖ 这个例子来自于 Treynor 和 Mazuy (1966), 一个早期对于共同基金择时能力的回归研究。为了刻画这种曲线关系,他们在模型中包含了超额市场收益率的平方项,而这并没有违背多元线性回归模型的假设,因为因变量和自变量关于系数是线性的。

信息。

9.5.2 函数形式误设定

每当我们对于回归进行估计时，我们必须假设回归具有正确的函数形式。这一假设可能会因为许多原因而无法满足：

- 一个或多个重要变量没有包含在回归中。
- 一个或多个回归变量可能需要在回归估计前进行转换（例如，对变量取自然对数）。
- 回归模型将本不应该放在一起的不同样本放在了一起。

首先，我们考虑回归中遗漏了一个重要的自变量所产生的影响（遗漏变量偏差）。如果真正的回归模型为

$$Y_i = b_0 + b_1X_{1i} + b_2X_{2i} + \epsilon_i \tag{9-7}$$

但是我们所估计的模型为[⊖]

$$Y_i = a_0 + a_1X_{1i} + \epsilon_i$$

那么我们的回归模型将是设定错误的。这样设定的模型的问题是什么呢？

如果遗漏的变量（ X_2 ）与仍在模型中的变量（ X_1 ）是相关的，那么模型的误差项就将与（ X_1 ）相关，于是回归系数的估计值 a_0 和 a_1 就会是有偏和不相合的。另外，这些系数标准误的估计值也将会是不相合的，所以我们既不能使用系数估计值，也不能使用估计出的标准误来进行统计检验。

例 9-10 遗漏变量偏差和买卖价差

在这个例子中，我们扩展对于买卖价差的检验，来给出回归中遗漏变量所产生的影响。在例 9-11 中，我们表明 [（买卖价差）/价格] 比率的自然对数显著地与做市商数量的自然对数和公司市值的自然对数相关。我们在下面再次给出例 9-1 中的表 9-1。

表 9-11 （重复的）ln（买卖价差/股价）和 ln（做市商数量）、ln（市值）的回归结果

			系数	标准误	t 统计量
截距			-0.758 6	0.136 9	-5.541 6
ln(纳斯达克做市商数量)			-0.279 0	0.067 3	-4.142 7
ln(公司市值)			-0.663 5	0.024 6	-27.008 7
ANOVA	df	SS	MSS	F	F 显著性
回归	2	2 681.648 2	1 340.824 1	1 088.832 5	0.00
残差	1 816	2 236.282 0	1 231 4		
总体	1 818	4 917.930 2			
残差标准误		1.109 7			
多元 R 平方		0.545 3			
观测数		1 819			

资料来源：FactSet, Nasdaq.

如果我们没有将市值的自然对数作为回归中的自变量包含在模型内，并将 [（买卖价差）/价格] 比率的自然对数仅与股票做市商数量的自然对数进行回归，那么我们会得到表 9-12 的结果。

⊖ 我们在 X_{2i} 被遗漏时采用一个不同的回归系数的记号，因为截距项和 X_{1i} 的斜率系数一般将与包含 X_{2i} 时所得到的结果不同。

表 9-12 ln（买卖价差/股价）和 ln（做市商数量）的回归结果

		系数	标准误	t 统计量	
截距		- 0. 122 9	0. 159 6	- 0. 769 8	
ln（纳斯达克做市商数量）		- 1. 662 9	0. 051 7	- 32. 151 9	
ANOVA	df	SS	MSS	F	F 显著性
回归	1	1. 783. 354 9	1 783. 354 9	1 033. 746 4	0. 00
残差	1 817	3 134. 5753	1 7251		
总体	1 818	4 917. 9302			
残差标准误		1. 313 4			
多元 R 平方		0. 3626			
观测数	1 819				

资料来源：FactSet，Nasdaq.

注意到 ln（纳斯达克做市商数量）前的系数从最初（正确设定的）回归中的 -0.279 0 下降到了错误设定的回归中的 -1.662 9。同样地，截距项从正确设定的模型中的 -0.758 6 上升到了错误设定的模型中的 -0.122 9。这些结果说明，遗漏一个本该在回归中的自变量可能造成余下变量前的回归系数不是相合的。

第 2 个造成回归模型错误设定的常见问题是对回归中的数据使用了错误的函数形式。而在这时，转化的数据形式才是恰当的。例如，有时分析师们无法通过所设定的变量间线性关系来解释因变量和一个或多个自变量间的曲线或者非线性关系。当我们设定一个回归模型时，我们应该考虑经济理论给出的是不是一个非线性关系。如下面的例 9-11 所示，我们经常可以通过将数据标在一张图上来确认这一非线性关系。如果当一个或者多个变量表示为变量的比例变动使得变量间的关系变成线性时，那么我们可以通过对我们想要表示为比例变动的变量取自然对数来修正错误设定。在其他的情况下，当标准化的数据（如净利润或者现金流除以销售量）更恰当时，分析师在回归中使用没有标准化的数据也会造成错误设定。在例 9-1 中，我们将买卖价差除以股价进行标准化，这是因为对于给定的投资规模，以交易成本表示的给定的买卖价差依赖于股票价格；如果我们没有将买卖价差进行标准化，那么回归就将是错误设定的。

例 9-11 买卖价差的非线性关系

在例 9-1 中，我们说明 [(买卖价差)/股价] 比率的自然对数显著地与做市商数量的自然对数和公司市值的自然对数有关系。但是为什么我们对回归中的每个变量取自然对数？我们在例 9-1 中开始的对于这个问题的讨论，现在我们将继续讨论下去。

那么，关于（买卖价差）/股价比率或者百分比买卖价差和其决定因素（自变量）之间关系本质的理论告诉我们什么呢？斯托尔（Stoll）（1978）建立了一个交易商市场中百分比买卖价差决定因素的理论模型。在他的模型中，决定因素是以乘积的特定形式进入模型的。以例 9-1 中所介绍的自变量来表示，他所假定的函数形式为

$$[(买卖价差)/股价]_i = c(做市商的数量)_i^{b_1} \times (市值)_i^{b_2}$$

式中的 c 是常数。在原始数据中，百分比买卖价差和做市商和市值的关系不是线性的。[⊖]然而，如果我们在上面模型的两边取自然对数，那么我们就得到一个关于变换变量线性的双对数回归模型：[⊗]

⊖ 模型的形式可以类比于经济学中柯布道格拉斯生产函数（Cobb-Douglas 生产函数）。
⊗ 我们已经在模型中加入了误差项。

$$Y_i = b_0 + b_1 X_{1i} + b_2 X_{2i} + \epsilon_i$$

式中， Y_i 表示第*i*只股票的（买卖价差）/股价比率的自然对数； b_0 表示等于 $\ln(c)$ 的自然对数； X_{1i} 表示第*i*只股票做市商数量的自然对数； X_{2i} 表示第*i*家公司市值的自然对数； ϵ_i 表示误差项。

正如例9-1中所提到的，双对数模型中的斜率系数被解释为弹性，更精确地说是因变量关于自变量的偏弹性（“偏”意味着保持其他自变量不变）。

我们可以通过将数据作图来判断在对数变换之后，变量间是否是线性相关的。例如，图9-4给出了一只股票做市商数量的自然对数（X轴）和（买卖价差）/股价比率的自然对数（Y轴）的散点图。另外，在图上画出了一条反映这两个变换后变量线性关系的回归直线。这两个变换后变量的关系显然是线性的。

如果我们没有对（买卖价差）/股价比率取对数，那么数据的图像将不是线性的。图9-5给出了一只股票做市商数量的自然对数（X轴）和（买卖价差）/股价比率（Y轴）的散点图。同样，在图上也画出了一条试图反映这两个变量间线性关系的回归直线。我们可以看到这两个变量间的关系显然是非线性的。^①结果，我们不应该对于以（买卖价差）/股价作为因变量的回归进行估计。为了保证预测的买卖价差为正，我们将不把（买卖价差）/股价作为因变量。如果我们使用没有经过变换的（买卖价差）/股价比率作为因变量，那么估计模型将可能得到负的买卖价差，这一结果将没有现实意义。事实上，买卖价差不可能为负（我们很难使得交易商同时高买低卖），所以我们会预测出负的买卖价差的模型一定是设定错误的。^②我们现在就来举例说明负的预测买卖价差的问题。

表9-13给出了（买卖价差）/股价作为因变量和做市商数量的自然对数以及公司市值的自然对数这两个自变量间的回归结果。

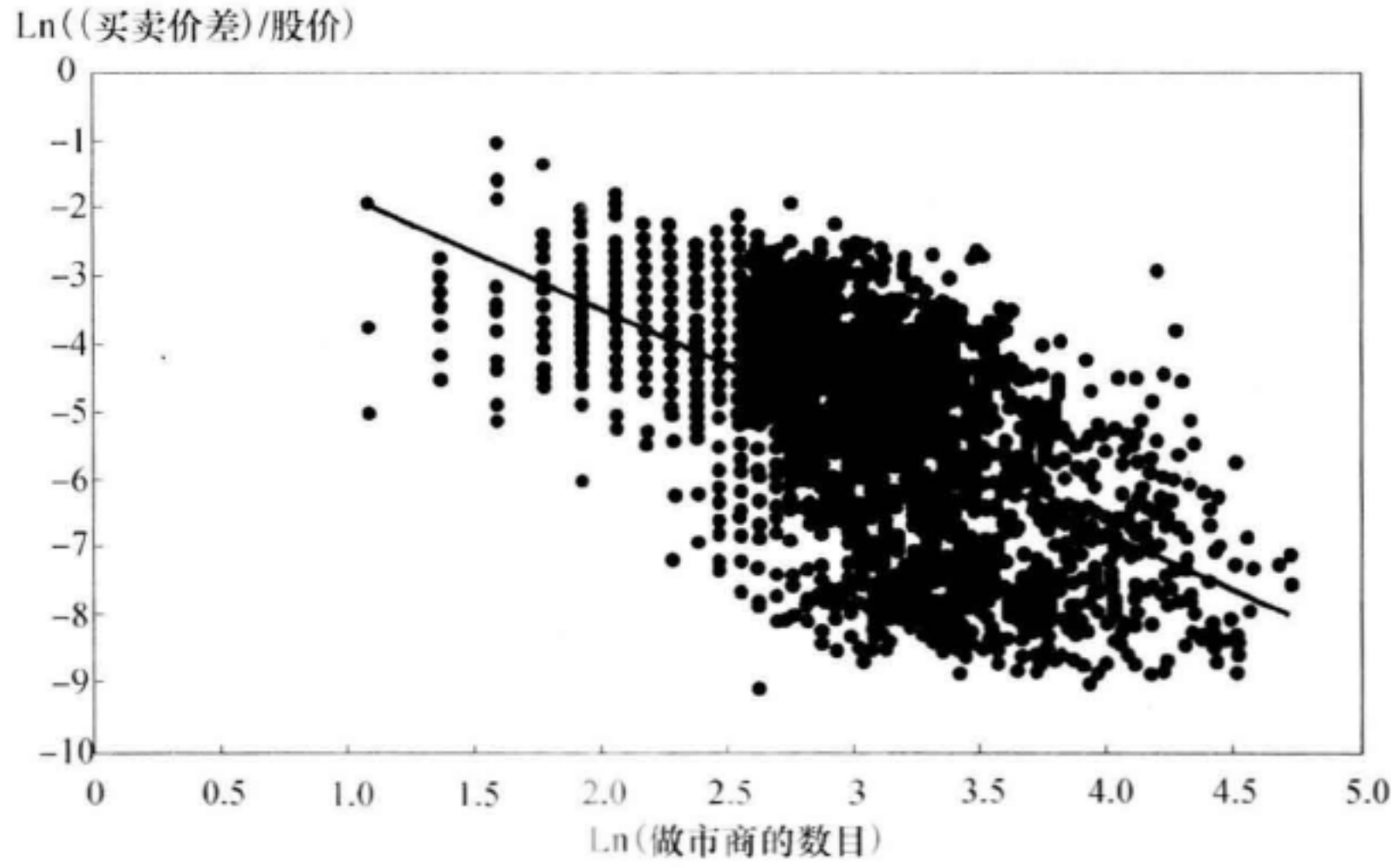


图9-4 当两个变量具有线性关系时的线性回归

① 虽然 $\ln(\text{买卖价差})/\text{股价}$ 和 $\ln(\text{市值})$ 之间的关系是线性的，但是，（买卖价差）/股价和 $\ln(\text{市值})$ 之间的关系是非线性的。我们为了节省篇幅，所以略去了这些散点图。
② 在我们的数据样本中，1819 家公司的每个买卖价差都是正的。

假设对于一只特定的在纳斯达克上市的股票，其做市商的数目是 20，而市值为 50 亿美元。因此，做市商数量的自然对数等于 $\ln 20 = 2.9957$ ，股票的市值（以百万美元计）为 $\ln 5\,000 = 8.5172$ 。在这些条件下，我们预测的（买卖价差）/股价的比率为 $0.0658 + (2.9957 \times (-0.0045)) + (-0.0068 \times 8.5172) = -0.0056$ 。因此，该模型预测（买卖价差）/股价的比率为 -0.0056 或者股价的 -0.56% 。因此，预测的买卖价差是负的，而这在经济上是没有道理的。我们可以通过使用（买卖价差）/股价的自然对数作为因变量来避免这个问题。[⊖]

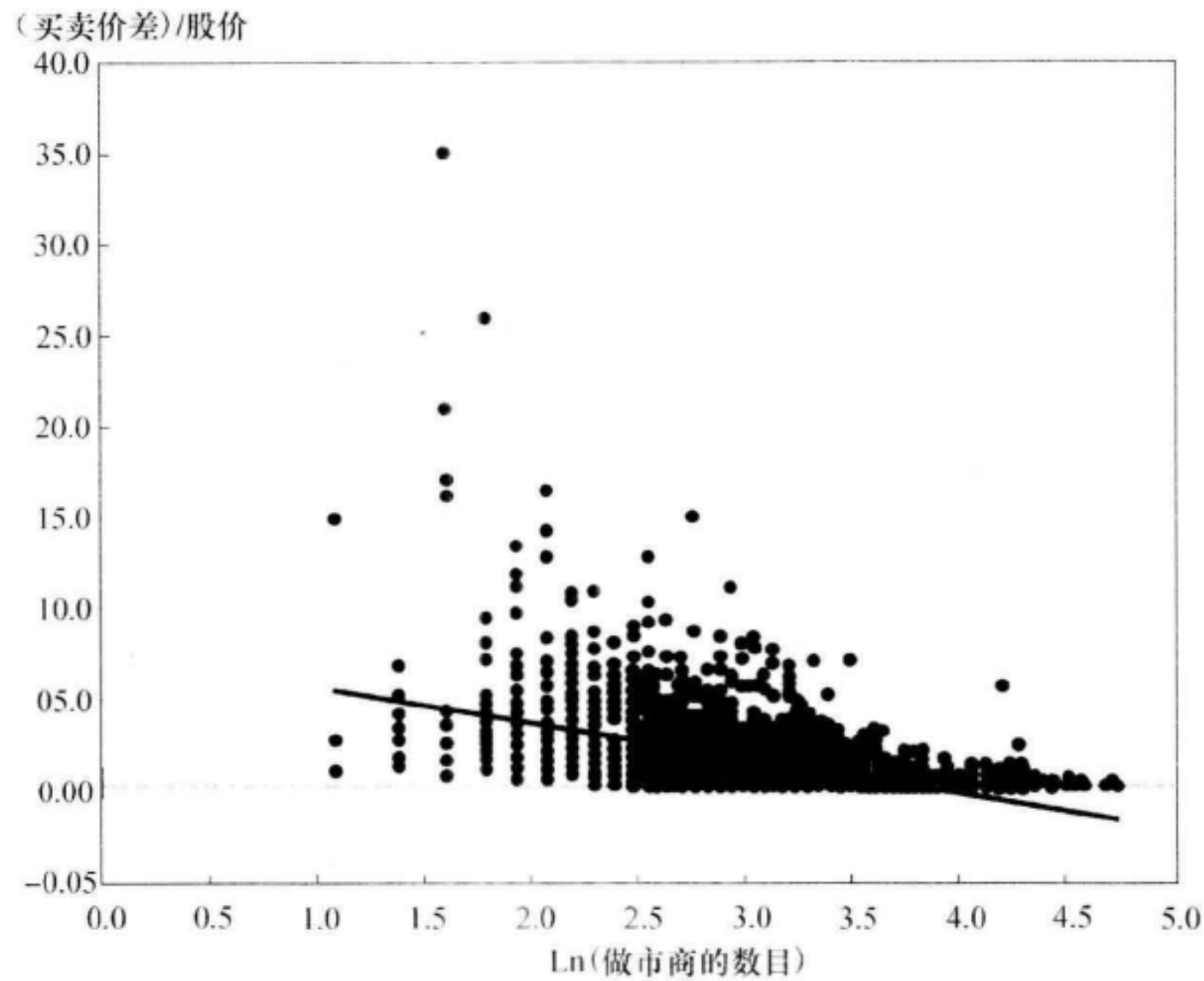


图 9-5 当两个变量具有非线性关系时的线性回归

表 9-13 （买卖价差）/股价关于 $\ln$ （做市商数量）和 $\ln$ （市值）回归的结果

			系数	标准误	t 统计量
截距			0.065 8	0.002 4	27.643 0
$\ln$ （纳斯达克做市商数量）			-0.004 5	0.001 2	-3.871 4
$\ln$ （公司市值）			-0.006 8	0.000 4	-15.867 9
ANOVA	df	SS	MSS	F	F 显著性
回归	2	0.318 5	0.159 2	427.817 4	0.00
残差	1 816	0.676 0	0.000 4		
总体	1 818	0.994 4			
残差标准误		0.019 3			
多元 R 平方		0.320 3			
观测数		1 819			

通常，分析师必须在他们比较公司间的数据前决定是否要调整数据。例如，在财务报告分析

⊖ 无论百分比买卖价差 Y 的自然对数为正还是为负，所求得的百分比买卖价差 e^Y 都是正的，因为一个正数的任何次方都是正的。常数 e 是正的 ($e \approx 2.7183$)。

中，分析师们经常使用同比财务报表（common size statements）来进行公司间的比较。在一个同比财务报表中，一家公司的利润表中的所有项都除以了该公司的营业收入。^①同比财务报表使得公司间的比较更加简单。分析师可以使用同比财务报表来快速比较一组公司净利润（或者利润表中的其他变量）的变化趋势。

可比性的问题同样也在分析师们使用回归分析对于一组公司的业绩进行比较时出现。例 9-12 对于这个问题进行了说明。

例 9-12 数据调整 and 经营现金流和自由现金流之间关系

假如一位分析师想要将一家公司的自由现金流解释为 2001 年美国 11 家在 2001 年年底市值超过 1 亿美元的家庭服装商店经营现金流的一个函数。

为了考察该问题，该分析师可能会使用以自由现金流作为因变量，而经营现金流作为自变量的单自变量线性回归。表 9-14 给出了这一回归的结果。注意到经营现金流前斜率系数的 t 统计量很高（6.5288），而回归的 F 统计量的显著性水平却很低（0.0001），并且 R 平方也很高。我们可能会认为该回归非常成功，所以对于该家庭服装店，如果经营现金流增长 1.00 美元，那么我们可能会自信地预测该公司的自由现金流会增加 0.3579 美元。

表 9-14 自由现金流关于家庭服装店经营现金流的回归结果

		系数	标准误	t 统计量	
截距		0.7295	27.7302	0.0263	
经营现金流		0.3579	0.0548	6.5288	
ANOVA	df	SS	MSS	F	F 显著性
回归	1	245 093.783 6	245 093.783 6	42.624 7	0.0001
残差	9	51 750.313 9	5 750.0349		
总体	10	296 844.097 5			
残差标准误		75.829 0			
多元 R 平方		0.8257			
观测数		11			

资料来源：Compustat.

但是，该模型的设定是否正确呢？该回归没有对于样本中公司规模的区别进行解释。

我们可以利用不同公司相同规模现金流的结果来解释公司规模差异。我们将经营现金流和自由现金流分别除以公司销售额调整后的变量用于回归分析。我们将（公司的自由现金流/销售额）作为因变量，而把（经营现金流/销售额）作为自变量。表 9-15 给出了这一回归的结果。注意到（经营现金流/销售额）前斜率系数的 t 统计量为 1.6262，所以它在 0.05 显著性水平下并不显著。另外，我们注意到 F 统计量的显著性水平为 0.1383，所以我们无法在 0.05 显著性水平下拒绝该回归无法解释（公司的自由现金流/销售额）在不同家庭服装店变动的原假设。最后，注意一下，该回归的 R 平方要远低于前面的回归。

① 关于共同比率财务报告的更多内容，参见 White, Sondhi 和 Fried（2003）。自由现金流和经营现金流在 Stowe, Robinson, Pinto 和 McLeavey（2002）中有具体的讨论。

表 9-15 家庭服装店的（自由现金流/销售额）关于（经营现金流/销售额）的回归结果

			系数	标准误	t 统计量
截距			-0.012 1	0.022 1	-0.549 7
经营现金流/销售额			0.474 9	0.029 20	1.626 2
ANOVA	df	SS	MSS	F	F 显著性
回归	1	0.003 0	0.003 0	2.644 7	0.138 3
残差	9	0.010 2	0.001 1		
总体	10	0.0131			
残差标准误		0.0336			
多元 R 平方		0.2271			
观测数		11			

资料来源：Compustat.

哪一个回归更有道理呢？一般，数据调整后的回归更有道理。我们想要了解如果经营现金流（作为销售额的一个百分比）发生一个变动，自由现金流（作为销售额的一个百分比）会发生什么变化。在没有考虑数据调整之前，回归的结果只是反映了不同公司计量上的差别，而不是反映公司数据背后所对应的经济学原因。

回归模型中第 3 种常见的错误设定为将本不应该放在一起的不同样本数据放在了一起。这类错误设定可以用图像来很好地反映。图 9-6 给出了两簇关于变量 X 和 Y 的数据，同时绘制了拟合的回归直线。这些数据可能反映了在公司成长的两个不同阶段的两个金融变量间的某种关系。

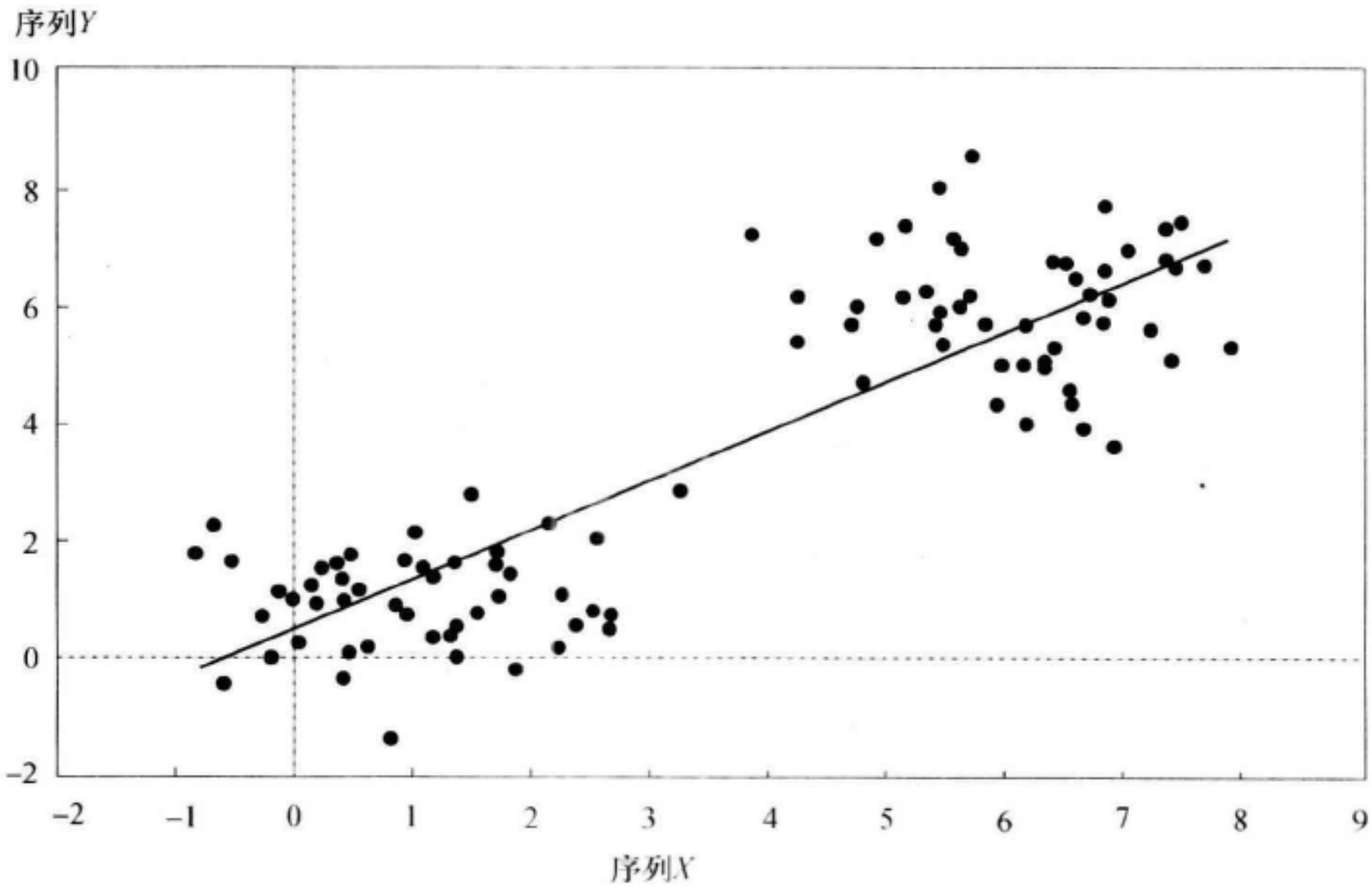


图 9-6 均值不同的两个序列的散点图

在每簇 X 和 Y 的数据中，这两个变量间的相关系数几乎为 0。因为 X 和 Y 的均值在联合样本中的两簇数据中是不同的，所以 X 和 Y 在联合样本中反映出很高的相关性。然而，该相关系数是虚假的（具有误导性的），因为它反映了公司在不同时期数据度量上的差别。

9.5.3 时间序列误设定 (自变量与误差相关)

在前一节中,我们讨论了当一个相关的自变量在回归中被遗漏时所产生的错误设定问题。在这一节中,我们将讨论由于不同种类的变量被包含在同一回归中,特别是在时间序列回归中,所可能产生的问题。在使用时间序列数据对于不同变量间关系进行解释的回归模型中,回归假设3 (即在给定自变量的条件下,误差项具有的均值为0) 特别容易被违背。如果这条假设被违背,那么估计的回归系数将是有偏和不相合的。

3 种常见的问题可能引起这种时间序列的模型设定错误:

- 回归中的自变量中包含了滞后因变量的项,使得回归具有序列相关的误差项;
- 回归中的自变量中包含了因变量的函数,有时,这一问题来自于变量在时间设定上的错误;
- 自变量的度量存在误差。

下面的例子对这些问题进行了说明。

假如一位分析师已经估计出了一个具有显著序列相关误差项的线性回归。这一序列相关性可以通过本章之前所讨论的方法进行修正。然而,假如分析师在回归中包含了因变量的一阶滞后项的另一个自变量。例如,分析师可能会使用回归方程

$$Y_t = b_0 + b_1X_{1t} + b_2Y_{t-1} + \epsilon_t \tag{9-8}$$

因为我们假设误差项是序列相关的,根据定义,误差项与因变量相关。结果,滞后的因变量 Y_{t-1} 将与误差项相关,从而违背了自变量与误差项不相关的假设。因此,回归系数的估计量将是有偏和不相合的。

例 9-13 带有因变量滞后项的费雪效应

在我们对于序列相关性的讨论中,我们利用杜宾 - 沃特森检验得到了费雪效应等式 (式 (9-5)) 中的误差项表现出正的相关性的结论。其中,我们使用 3 个月的国库券收益率作为因变量,而用专业预测家对于通货膨胀率的预期作为自变量。因变量和自变量的观测值使用的是季度数据。表 9-16 给出了修改后的回归结果。在回归中我们在自变量中加入了前一季度的 3 个月国库券的收益率。

表 9-16 国库券收益率关于预期通货膨胀率和国库券收益率滞后项的回归结果

	系数	标准误	t 统计量
截距	0.004 6	0.004 0	1.571 8
预期通货膨胀率	0.275 3	0.063 1	4.361 0
国库券收益率的滞后项	0.755 3	0.049 5	15.251 0
残差标准误		0.013 4	
多元 R 平方		0.804 1	
观测数		137	

资料来源: Federal Reserve Bank of Philadelphia, U. S. Department of Commerce.

乍看之下,这些回归结果看上去非常有趣,国库券收益率的滞后项前的系数似乎很显著。但是在进一步考察之后,我们必须忽视这些回归结果,因为回归模型在本质上是设定错误的。只要误差项是序列相关的,在回归中的自变量中包含了国库券收益率的滞后项将会导致所有的系数估计值有偏和不相合。因此,这个回归结果既不能用来假设检验,也不能用来预测。

第2种在投资分析中常见的有关时间序列的错误设定为预测过去。这是什么意思呢？如果我们预测将来（比如说我们在时刻 t 对于时刻 $t+1$ 的变量 Y 的值进行预测），我们必须基于我们在时刻 t 所知道的信息进行预测。我们可能会使用下面的等式进行回归来做出预测

$$Y_{t+1} = b_0 + b_1 X_{1t} + \epsilon_{t+1} \tag{9-9}$$

在这个等式中，我们使用时刻 t 的 X 值来预测时刻 $t+1$ 的 Y 值。误差项 ϵ_{t+1} 是在时刻 t 未知的，因此其应该与 X_{1t} 不相关。

而不幸的是，分析师们有时会使用这样的回归模型，该模型使用时刻 $t+1$ 因变量值的函数来试图预测时刻 $t+1$ 的因变量的值。在这种模型中，自变量将与误差项相关，因此，该回归方程是设定错误的。举一个例子，一位分析师可能试图利用一组公司的市值账面比和其年底的市值来解释该年这些公司的横截面收益率。[⊖]如果分析师认为这一回归能够对高市值账面比或者高市值的公司是否会具有更高的收益率做出预测的话，那么该分析师肯定会得到错误的结论。对于任一给定区间，该区间中的收益率越高，在期末的市值也就越高。而对于市值账面比也同样如此（即区间中的收益率越高，市值账面比也就越高）。所以在这个例子中，如果所有的横截面数据来自于第 $t+1$ 个区间，那么因变量（收益率）较高的值实际上直接导致了自变量（市值和市值账面比）较高的值，而不是其他的结果。在这种类型的错误设定中，回归模型实际上在回归等式左右两边都包含了因变量的值。

第3种常见的有关时间序列的错误设定是在自变量存在度量误差时产生的。假如一个金融理论告诉我们某个变量 X_t ，如预期通货膨胀率，应该被包含在回归模型中。而我们无法观测到 X_t ；而作为替代，我们观测到实际通货膨胀率， $Z_t = X_t + u_t$ ，其中 u_t 为一个和 X_t 不相关的误差项。即使在这么好的设定之下，在回归中使用 Z_t 而不是 X_t 同样也会造成回归系数的估计值有偏和不相合。下面让我们来看一下为什么是这样的。如果我们想要估计回归式

$$Y_t = b_0 + b_1 X_t + \epsilon_t$$

但是我们观测到 Z_t 而不是 X_t ，于是我们转而估计

$$Y_t = b_0 + b_1 Z_t + (-b_1 u_t + \epsilon_t)$$

因为 $Z_t = X_t + u_t$ ，所以 Z_t 与误差项 $(-b_1 u_t + \epsilon_t)$ 是相关的。因此，我们的估计模型违背了误差项必须与自变量不相关的假设。结果，回归系数的估计值将是有偏和不相合的。

例 9-14 存在度量误差的费雪效应

回忆一下例 9-8 中对于费雪效应的检验。在那个例子中，我们无法拒绝 3 个月国库券收益率和预期通货膨胀率同步变动的原假设。

表 9-16 （重复的）国库券收益率和预期通货膨胀率的回归结果

	系数	标准误	t 统计量
截距	0.0304	0.0040	7.6887
预期通货膨胀率	0.8774	0.0812	10.8096
残差标准误		0.0220	
多元 R 平方		0.4640	
观测数		137	
杜宾 - 沃特森统计量		0.4673	

资料来源：Federal Reserve Bank of Philadelphia, U. S. Department of Commerce.

⊖ “市值账面比”是每股价格除以每股账面价值后得到的比值。

如果我们使用实际通胀率而不是预期通货膨胀率作为自变量，那么会发生什么呢？首先我们要注意

$$\pi = \pi^e + v$$

式中， π 表示实际通货膨胀率； π^e 表示预期通货膨胀率； v 表示实际通货膨胀率和预期通货膨胀率之间的差别。

因为实际通货膨胀率反映了包含误差的预期通货膨胀率，所以使用国库券收益率作为因变量，实际通货膨胀率作为自变量的回归中所得到的回归系数的估计值将不是相合的。[Ⓐ]

表 9-17 给出了使用实际通货膨胀率作为自变量的回归结果。在这个表中的估计值与我们在之前表中所给出的结果存在很大不同。我们注意到实际通货膨胀率前的斜率系数远低于前一个回归中预期通货膨胀率前的斜率系数。这个结果表明这样一个一般性的结论：在单自变量的回归中，如果我们选择一个存在度量误差的自变量，那么那个变量前的斜率系数的估计值将是有偏的，而偏向于 0。[Ⓑ]

表 9-17 国库券收益率和实际通货膨胀率的回归结果

	系数	标准误	t 统计量
截距	0.0432	0.0034	12.7340
实际通货膨胀率	0.5066	0.0556	9.1151
残差标准误		0.0237	
多元 R 平方		0.3810	
观测数		137	

资料来源：Federal Reserve Bank of Philadelphia, U. S. Department of Commerce.

9.5.4 其他类型时间序列误设定

目前，在线性回归中最常见的错误设定的原因是使用两个或多个不平稳的时间序列变量。粗略地说，非平稳性（nonstationarity）表示一个变量的特征，如均值和方差，不是随着时间保持不变的。我们将对于平稳性的讨论推迟到第 10 章中进行，但是我们可以在使用回归中的统计推断之前，列出一些我们需要进行平稳性检验的例子。[Ⓒ]

- 具有趋势的时间序列间的关系（例如，消费和 GDP 之间的关系）。
- 可能是随机游走（random walk）的时间序列间的关系（下一期最佳的预测值是本期的观测值的时间序列）。汇率经常是随机游走的时间序列。

本章中时间序列的例子是经过我们仔细挑选的，它们都不可能存在不平稳的问题。但是非平稳性可能在分析两个或者两个以上时间序列间关系时是一个非常严重的问题。分析师们必须在使用线性回归对这些时间序列间的关系进行分析之前，了解到这些问题。不然，他们可能得到

Ⓐ 相合估计量是指这样一种估计量：该估计量随着样本容量的增大，其值接近真实的总体参数值的概率也会增加，并趋近于 1。
Ⓑ 这个命题并不适用于多于一个自变量的回归。当然，我们在这个例子中忽略了序列相关的误差，但是因为回归系数是不相合的（因为有度量误差），所以在这里对于序列相关的检验和修正是毫无意义的。
Ⓒ 术语“平稳性检验”中包含单位根检验和协整检验。这些检验将在第 10 章中进行讨论。

无效的统计推断结果。

9.6 因变量是定性变量的模型

金融分析师们经常需要有能力对于定性因变量的结果进行解释。定性因变量 (qualitative dependent variables) 是用作因变量而非自变量的虚拟变量。

例如,为了预测一家公司是否会破产,我们需要使用定性因变量(破产与否)作为因变量,同时使用公司财务业绩表现的数据(如权益收益率、负债权益比或者负债评级)作为自变量。不幸的是,线性回归并不是最好的用来估计这种模型的统计模型。如果我们使用定性因变量(破产(等于1)或者不破产(等于0))作为回归中的因变量,而使用财务变量作为自变量,那么,因变量的预测值可能会比1更大或者比0更小。当然,这些结果是没有意义的。破产的概率(或者类似的其他概率值)不可能超过100%或者低于0%。作为线性回归模型的替代,我们应该对于这种估计问题使用probit模型或者判别分析。

Probit和Logit模型(probit and Logit models)在给定所用来解释因变量结果的自变量值的条件下给出离散结果的概率估计值。基于正态分布的Probit模型在给定自变量 X 值的条件下,估计出 $Y=1$ (条件得到满足)概率的估计值。Logit模型是类似的,区别仅仅在于其基于的是逻辑斯谛分布而不是正态分布。[⊖]这两个模型必须使用最大似然方法进行估计。[⊗]

另一个处理定性因变量的技术是判别分析(discriminant analysis)。在奥尔特曼(Altman)(1968,1977)的 Z 得分(Z -score)和西塔分析(Zeta analysis)中,他给出了判别分析的结果。奥尔特曼使用财务比率来预测表示破产的定性因变量。判别分析使用一个线性函数,类似于一个回归方程,来求得一个总体的得分。基于这个得分,一个观测值可以被分成破产或者没有破产两种类型。

不仅在投资组合管理中,而且在商业管理中,定性因变量模型都是很有用的。例如,我们可能需要预测一位客户是否可能继续投资于一家公司或者从该公司中收回其投资资产。我们也可能想要解释某个人口统计学的特征是如何影响一个潜在投资者签约成为一个新客户的概率,或者想要基于目标受众的人口统计学特征来对于某个直接发送邮件的广告运动的效果进行评价。这些问题都可以使用probit或者logit模型进行分析。

例 9-15 对于分析师荐股的解释

假如我们想要研究什么因素决定了是否至少有一位分析师推荐一家公司的结果出现。我们使用一个probit模型来回答这个问题。我们的样本包括1999年上市公司的2047个观测数据。所有数据来自于Disclosure, Inc.。关于Disclosure的分析师荐股数据来自于I/B/E/S。

Probit模型中的变量如下:

ANALYSTS表示一个离散的因变量。如果至少有一位分析师推荐了这家公司,其值取1。而如果没有分析师推荐,其值取0。

⊖ 逻辑斯谛分布 $e^{(b_0+b_1X)} / [1 + e^{(b_0+b_1X)}]$ 比累积正态分布更容易计算。因此,logit模型在计算成本非常昂贵时更受欢迎。

⊗ 更多关于probit和logit模型的内容,参见Greene(2003)。

LNOLUME 表示最近一周交易量的自然对数。
LNMV 表示市场价值的自然对数。
ESTABLISHED 表示一个虚拟自变量。如果公司的财务数据在最近 5 年审计过的话,该值取 1。
LNTA 表示总资产(账面值)的自然对数。
LNSALES 表示净销售额的自然对数。

在试图对于分析师荐股的解释中,我们预计表示市场(交易量和价值)和账面(价值和销售额)的变量可能从不同的规模和重要性角度来对于分析师荐股进行解释。[⊖]有关审计历史的变量反映了一个可能的合适水平,在该水平下我们预计分析师可能获得审计过的财务报告。该模型包含 3 个我们认为可能与分析师荐股相关的变量(LNMV, LNTA 和 LNSALES)。

基于没有在这里给出的分析,我们的 probit 回归没有显示出多重共线性的经典症状。表 9-18 给出了 probit 回归的估计结果。

表 9-18 利用 probit 模型来解释分析师荐股

	系数	标准误	t 统计量
截距	-7.973 8	0.436 2	-18.281 5
LNOLUME	0.157 4	0.015 8	9.948 2
LNMV	0.444 2	0.036 9	12.026 8
ESTABLISHED	0.316 8	0.104 5	3.032 0
LNTA	0.054 8	0.0296	1.849 4
LNSALES	0.050 7	0.026 6	1.905 9
预测正确的百分比		73.67	

资料来源: Disclosure, Inc.

正如表 9-18 所示,(除了截距之外的)3 个系数的 t 统计量的绝对值都大于 2.0。LNOLUME 前系数的 t 统计量为 9.95。该值远高于 0.05 显著性水平下 t 统计量的临界值(1.96),所以我们能在 0.05 显著性水平下拒绝 LNOLUME 前的系数等于 0 的原假设,并接受系数不等于 0 的备择假设。第 2 个系数的 t 检验量绝对值大于 2.0 的变量是 LNMV,其 t 统计量为 12.03。我们也能在 0.05 显著性水平下拒绝 LNMV 前的系数等于 0 的原假设,并接受系数不等于 0 的备择假设。最后,ESTABLISHED 前系数的 t 统计量为 3.03。我们能在 0.05 显著性水平下拒绝 ESTABLISHED 前的系数等于 0 的原假设。

在 probit 分析中剩下的两个自变量在 0.05 显著性水平下都在统计上不显著。这两个变量前系数的 t 统计量的绝对值都没有超过 1.91,所以没有一个数达到拒绝原假设(相对应的系数显著不同于 0)所需的临界值 1.96。这一结果表明,一旦我们考虑了公司市场价值、交易量以及是否有 5 年的审计历史之后,其他因素(资产账面价值和销售价值)对于是否至少有一位分析师会推荐该公司是没有解释力的。

⊖ 更多关于多重共线性检验的内容,参见 Greene (2003)。

第10章

时间序列分析

10.1 引言

作为金融分析师，我们经常会使用时间序列数据来进行投资决策。一个时间序列（time series）是指一组关于某个变量在不同时段所出现结果的观测值序列：比如，过去五年中某家公司的季度销售额或者某只交易股票的日回报率。在本章中，我们将探讨时间序列模型的两种主要应用：解释过去和预测未来。我们也将讨论如何对于时间序列模型进行估计，并考察一个描述某个时间序列的模型是如何随时间而变化的。下面的两个例子列举出了我们可能想要对于时间序列所要提出的各种问题。

假如现在是2003年年初，我们正管理着一个由美国投资公司所持有的投资组合（该组合内包含瑞士的股票）。因为该投资组合的价值会在其他因素保持不变的情况下，随着瑞士法郎相对于美元的贬值而下降（反之亦然），所以我们现在考虑是否要对该投资组合所面临的法郎价值变化的风险暴露进行对冲。为了帮助我们做出这个决策，我们决定对于法郎/美元汇率的时间序列进行建模。图10-1绘制出了法郎/美元汇率月度数据的变化情况。（该数据用的是各个月度中日汇率的平均值。）我们可能会问：自1987年以来的汇率是否比1987年之前更加稳定？汇率是否表现出一个长期的趋势？我们如何最好地利用历史的汇率来预测未来的汇率？

作为另一个例子，假设现在是2001年年初。我们为一家卖方公司对零售商店进行研究，并想要预测下一年的零售额情况。图10-2给出了美国实际零售额的月度数据。这些数据都已经对通货膨胀进行了调整，但还未对季节性效应进行调整，因此，在每年岁末年初的节假日周围呈现出尖峰的形态。因为在零售商店财务报告中的销售额并没有进行季节性调整，所以我们将对于没有进行季节调整的零售额进行建模。那么，我们如何对于零售额中的趋势进行建模呢？我们如何对于周期性区间中发生的高峰和低谷的极端季节性效应进行调整呢？我们如何利用过去的零售额来预测未来的零售额呢？这些都是我们可能提出的问题。

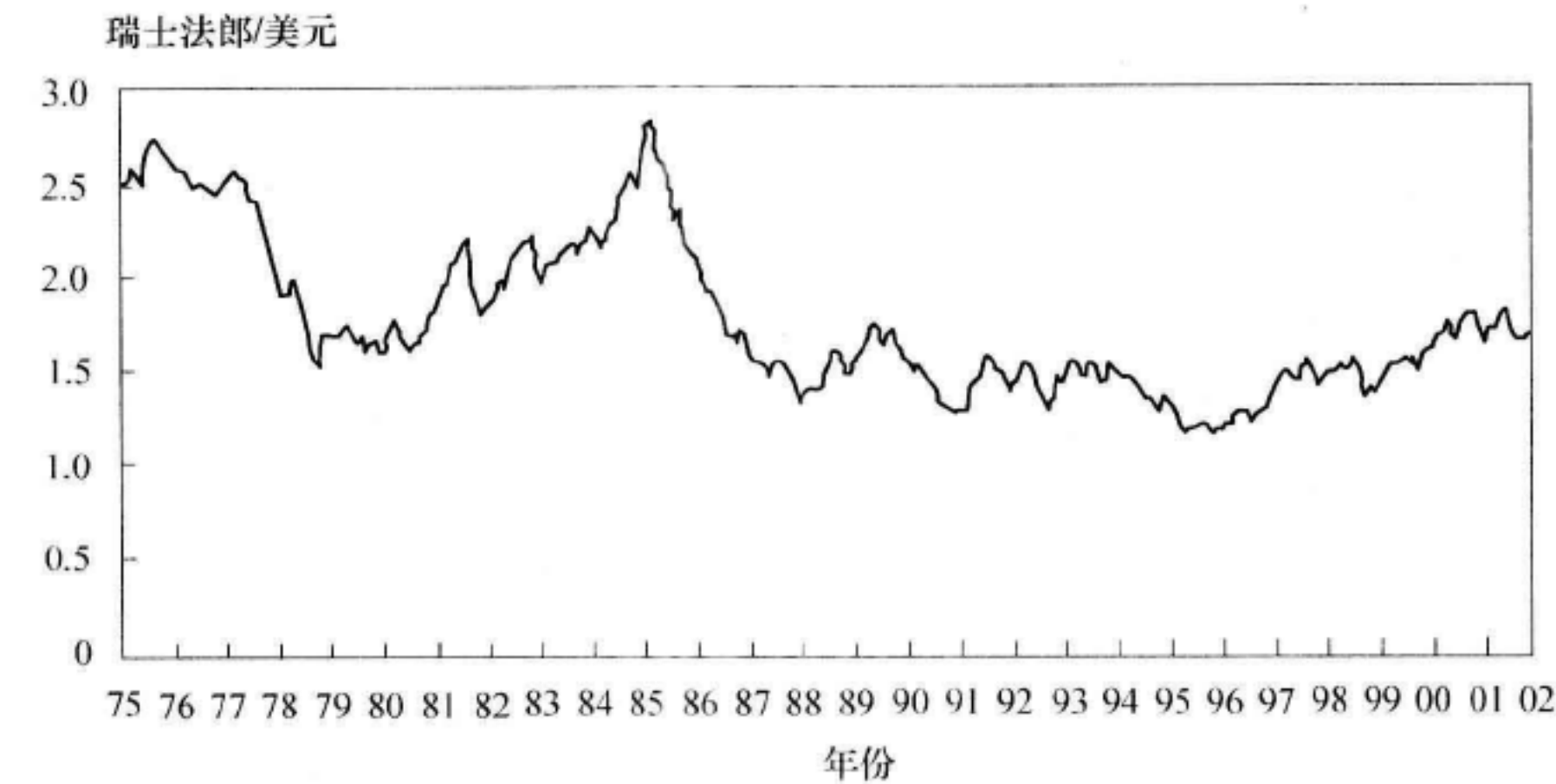


图 10-1 瑞士法郎/美元汇率，月度日汇率的平均值

资料来源：Board of Governors of the Federal Reserve System.

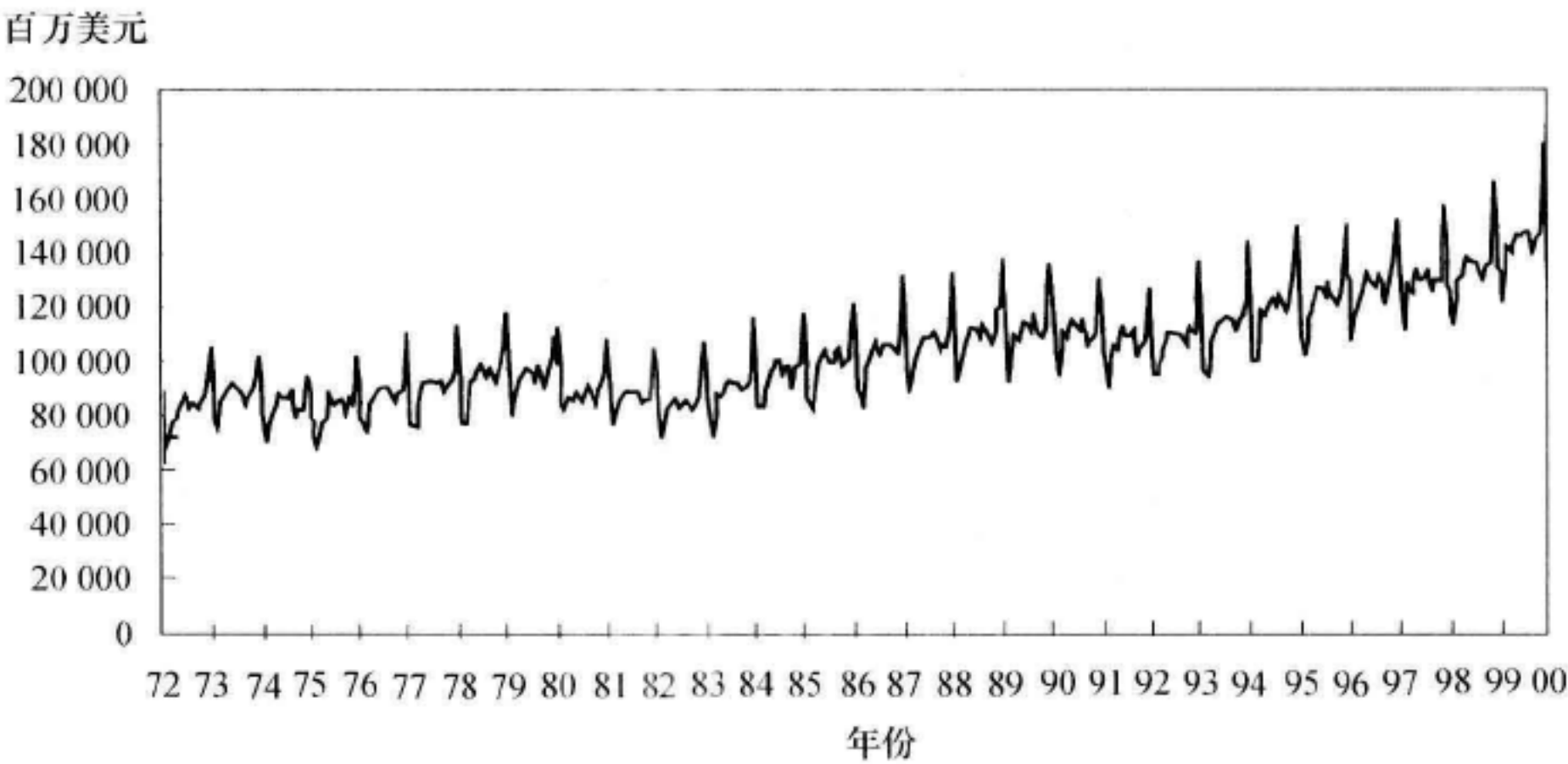


图 10-2 美国月度实际零售额

资料来源：U. S. Department of Commerce, Census Bureau.

时间序列分析中的一些基本问题为：我们如何对于趋势建模？如何基于历史数值对于未来值进行预测？如何对于季节性效应建模？如何在不同的时间序列模型中进行选择？以及如何对于时间序列中的时变方差进行建模？我们将在本章中一一对于这些问题予以解答。

10.2 处理时间序列数据所面临的挑战

在本章中，我们所要达到的目标是将线性回归的方法应用于一个给定的时间序列之中。而不幸的是，在对于时间序列进行建模时，我们经常发现线性回归模型的假设是无法得到满足的。为了使我们能够进行时间序列分析，我们需要确保线性回归模型的假设得到满足。当那些假设没有得到满足时，在许多情况下，我们会对时间序列进行变换，或者设定不同的回归模型，以使得线性回归模型的假设仍旧在新的设定下得到满足。

我们可以在一个常见的时间序列模型（自回归模型）的背景下，说明满足线性回归模型的假设所面临的困难。通俗地说，一个自回归模型是指自变量是因变量滞后值（即历史值）的一类模

型, 例如 $x_t = b_0 + b_1 x_{t-1} + \epsilon_t$ 。[⊖] 我们经常在处理时间序列时遇到如下问题:

- 残差的误差是相关的, 而不是不相关的。

在所计算的回归中, x_t 和 $b_0 + b_1 x_{t-1}$ 之间的差别被称为残差的误差 (residual error)。线性回归模型中假设该误差项在不同观测值间是互不相关的。在将时间序列的历史值作为自变量的时间序列模型中, 对于该假设的违背经常会比其他的因变量和自变量是不同的模型 (如横截面模型) 产生更加严重的后果。正如我们在多元回归一章中所讨论的, 在因变量和自变量是不同的回归中, 模型中误差的序列相关并不会影响截距或斜率系数估计值的相合性。而相反, 在一个自回归时间序列模型中 (如 $x_t = b_0 + b_1 x_{t-1} + \epsilon_t$), 误差项的序列相关性将会造成截距 (b_0) 和斜率系数 (b_1) 估计值的不相合。

- 时间序列的均值和 (或) 方差是随时间变化的。

如果我们对于一个具有时变均值和方差的时间序列用自回归模型进行估计, 那么回归结果将是无效的。

在我们试图使用时间序列模型进行预测之前, 我们可能需要转换时间序列模型以使得其能够满足线性回归的设定。在了解到该目标之后, 我们将会发现时间序列分析是相对直接并且富有逻辑的。

10.3 趋势模型

估计一个时间序列的趋势, 并利用该趋势来对于时间序列未来值进行预测是最简单的一种预测方法。例如, 我们在图 10-2 中所看到的, 月度美国实际零售额呈现出一个长期向上运动的形态——即一个趋势 (trend)。在这一节中, 我们将考察两种趋势, 线性趋势和对数线性趋势, 并讨论如何在它们之间做出选择。

10.3.1 线性趋势模型

最简单的一类趋势是线性趋势 (linear trend), 一种因变量以固定的速率随时间变动的趋势。如果一个时间序列 y_t 具有线性趋势, 那么我们可以利用下面的回归方程来对于该序列进行建模:

$$y_t = b_0 + b_1 t + \epsilon_t, t = 1, 2, \dots, T \quad (10-1)$$

式中, y_t 表示 t 时刻时间序列的值 (因变量的值); b_0 表示 y 的截距项; b_1 表示斜率系数; t 表示时间, 因变量或者解释变量; ϵ_t 表示一个随机误差项。

在式 (10-1) 中, 趋势线 $b_0 + b_1 t$ 给出了在时刻 t 时间序列的预测值 (其中的 t 在样本第一期取值为 1, 然后在之后每期依次在前一期的基础上增加 1)。因为系数 b_1 是趋势线的斜率, 所以我们将 b_1 称为趋势系数。我们可以利用普通最小二乘法估计出 b_0 和 b_1 这两个系数, 并将其估计出的系数分别记为 $\hat{b}_0$ 和 $\hat{b}_1$ 。[⊖]

现在我们将介绍如何使用这些估计量来对于一个特定时间段的时间序列值进行预测。我们记得 t 在第一期时取值为 1。因此, 在第一期 y_t 的预测值或者拟合值为 $\hat{y}_1 = \hat{b}_0 + \hat{b}_1(1)$ 。同样地, 在接下来的几期, 比如第六期, 拟合值为 $\hat{y}_6 = \hat{b}_0 + \hat{b}_1(6)$ 。现在假设我们想要预测在样本外某个时间段的时间序列值, 比如时段 $T+1$ 。那么, 时段 $T+1$ 的 y_t 的预测值为 $\hat{y}_{T+1} = \hat{b}_0 + \hat{b}_1(T+1)$ 。例如, 如果 $\hat{b}_0$ 为 5.1, 而

⊖ 我们也可以将等式写为 $y_t = b_0 + b_1 y_{t-1} + \epsilon_t$ 。

⊖ 回忆一下, 普通最小二乘法是一种基于最小化回归平方残差和准则的估计方法。

为 2，那么在 $t=5$ 处的 y_5 的预测值为 15.1，而在 $t=6$ 处的 y_6 的预测值为 17.1。我们注意到无论前一个时段的时间序列值处在什么水平，时间序列中每个之后的观测值都会比前一个增加 $\hat{b}_1=2$ 。

例 10-1 美国消费者价格指数的趋势

现在是 2001 年 1 月。作为一家银行信托部门的固定收益分析师，莉塞特·米勒 (Lisette Miller) 正在关注未来通货膨胀的水平以及其可能会对投资组合价值所产生的影响。为此，她想要预测未来的通货膨胀率。为了达到该目的，她首先需要估计出通货膨胀率中的线性趋势。于是，她利用美国月度消费者价格指数 (Consumer Price Index, CPI) 通货膨胀率的数据，该数据是以年度百分率来表示的，[⊖]其具体情况反映在图 10-3 中。该数据包含从 1985 年 1 月至 2000 年 12 月的 192 个月份的数据，而她所要估计的模型为 $y_t = b_0 + b_1t + \epsilon_t, t = 1, 2, \dots, 192$ 。表 10-1 给出了该等式的估计结果。由于具有 192 个观测值和 2 个参数，所以这个模型的自由度为 190。在 0.05 的显著性水平下， t 统计量的临界值为 1.97。因为系数 t 统计量的绝对值都远高于临界值，所以截距 ($\hat{b}_0=4.1342$) 和趋势系数 ($\hat{b}_1=-0.0095$) 都在统计上显著。于是，所估计的回归等式可以写为

$$y_t = 4.1342 - 0.0095t$$

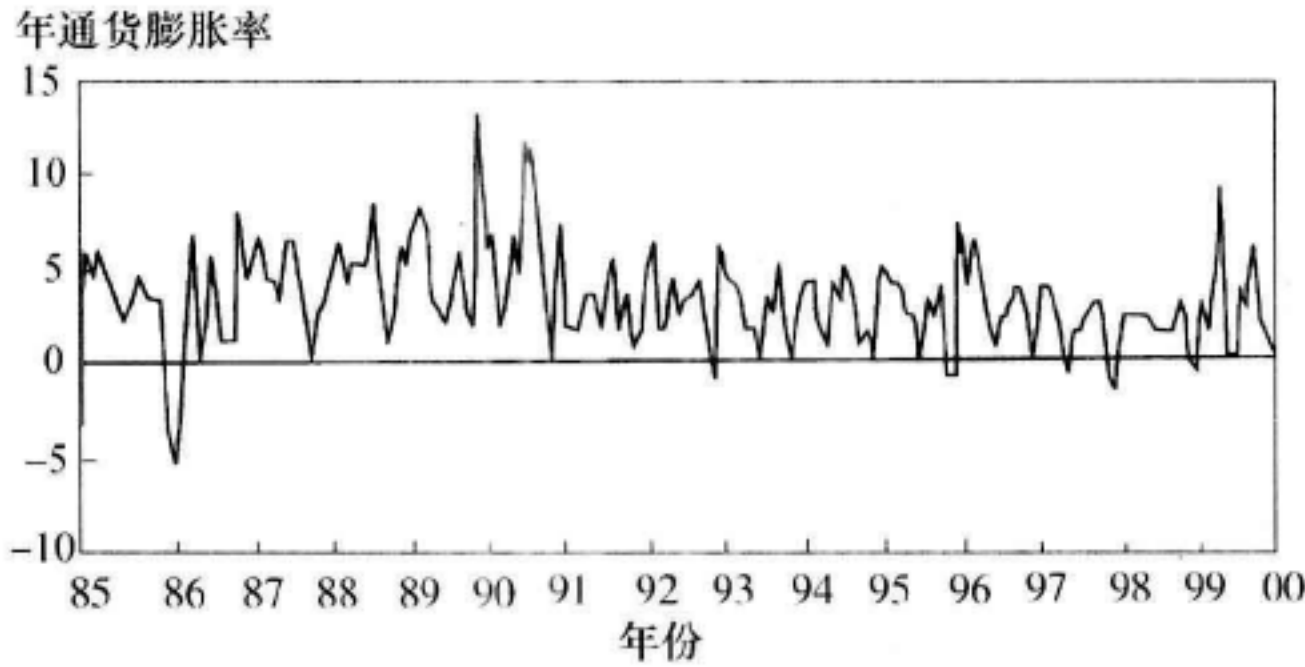


图 10-3 没有经过季节调整的月度 CPI 通货膨胀率

表 10-1 1985 年 1 月~2000 年 12 月通货膨胀月度观测值线性趋势的估计结果

回归统计量			
	R 平方	0.040 8	
	标准误	2.554 4	
	观测数	192	
	杜宾-沃特森	1.38	
	系数	标准误	t 统计量
截距	4.134 2	0.370 1	11.169 3
趋势	-0.009 5	0.003 3	-2.844 5

资料来源：U. S. Bureau of Labor Statistics.

因为趋势线斜率的估计值为 -0.009 5，所以米勒认为线性趋势模型在样本区间内的最佳估计为年度通货膨胀率大约以每月 1% 的速率下降。

⊖ 在这些数据中，1% 表示为 1.0。

在1985年1月,即样本的第一个月份,通货膨胀率的预测值为 $\hat{y}_1 = 4.1342 - 0.0095(1) = 4.1247(\%)$ 。而在2000年12月,第192个月份或者样本的最后一个月份,通货膨胀率的预测值为 $\hat{y}_{192} = 4.1342 - 0.0095(192) = 2.3177(\%)$ 。^①然而,我们注意到这些预测值都位于样本区间内。若将这些值与实际值相比较,我们就可以看出米勒模型拟合数据的好坏情况;然而,该估计模型的一个更加主要的目的是对于样本外区间的通货膨胀率水平进行预测。例如,对于2001年12月(样本最后一期的12个月之后), $t = 192 + 12 = 204$,此时,通货膨胀水平的预测值为 $\hat{y}_{204} = 4.1342 - 0.0095(204) = 2.2041(\%)$ 。

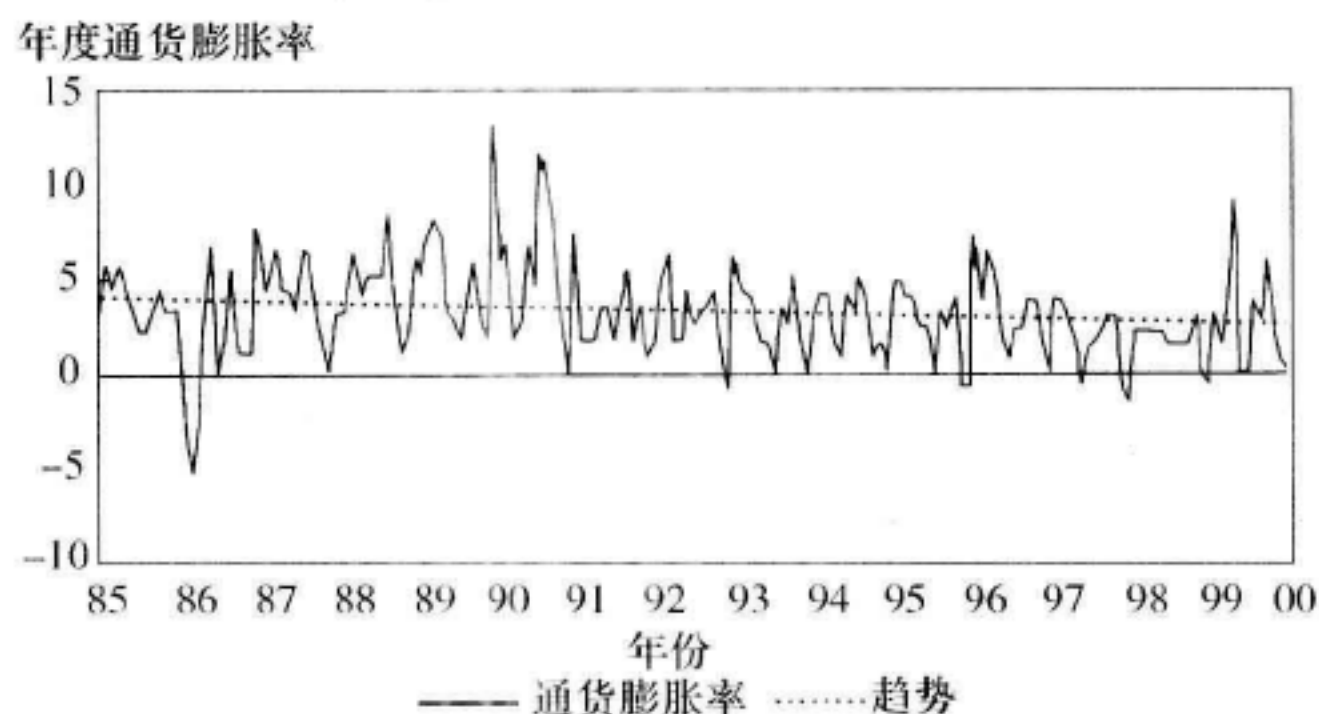


图 10-4 具有趋势的月度 CPI 通货膨胀率

资料来源: U. S. Bureau of Labor Statistics.

图 10-4 给出了通货膨胀率数据及其拟合出的趋势。注意到通货膨胀率没有在较长时间段表现出高于或者低于趋势线的情况。在趋势和实际通货膨胀率间没有持续的偏差存在。残差(实际数据减去趋势值)表现出时间上的不可预测性和不相关性。因此,利用线性趋势线对于1985~2000年的通货膨胀率进行建模是合理的。而且,我们也可以得到通货膨胀率在那个时间段稳步下降的结论。同时注意到该模型的 R^2 很低,这意味着利用该模型对于通货膨胀进行预测的较大不确定性。事实上,趋势仅仅解释了月度通货膨胀变动的4.08%。在本章之后的部分,我们将对于我们是否可以对通货膨胀建立一个比仅仅使用趋势线更好的模型的问题进行考察。

10.3.2 对数线性趋势模型

有时,线性趋势不能准确地对于一个时间序列的增长率建模。在这种情况下,我们通常会发现若利用一个线性趋势对该时间序列进行拟合就会导致持续相关的而不是不相关的误差。如果一个线性趋势模型的残差是持续相关的,那么我们就需要使用其他的满足线性回归条件的模型。对于金融时间序列,一个重要的线性趋势的替代模型是对数线性趋势模型。对数线性趋势对于以指数增长的时间序列拟合得较好。

指数增长意味着以某个特定速率恒定增长。所以,以5%的恒定速率的年增长率是指数的,因为序列会向上不断增长到无限。那么,指数增长是如何进行的呢?假设我们用下面的等式来描述一个时间序列:

$$y_t = e^{b_0 + b_1 t}, t = 1, 2, \dots, T \quad (10-2)$$

那么,指数增长就是以—个恒定的速率($e^{b_1} - 1$)连续复利进行增长。例如,考虑时间序列在两

① 在报告最终结果(这里是2.3177)时,我们使用没有经过四舍五入的回归系数的估计值;而在给出的计算中,我们使用了经过四舍五入的回归系数,所以使用经过四舍五入的系数进行计算经常造成稍稍不同的结果。

个相邻时间段的值。在第一期时间序列的值为 $y_1 = e^{b_0 + b_1(1)}$ ，而在第二期，该值为 $y_2 = e^{b_0 + b_1(2)}$ 。前两期时间序列值的比率为 $y_2/y_1 = (e^{b_0 + b_1(2)}) / (e^{b_0 + b_1(1)}) = e^{b_1}$ 。一般地，在任一时段 t ，时间序列的值为 $y_t = e^{b_0 + b_1(t)}$ 。在时段 $t+1$ ，时间序列的值为 $y_{t+1} = e^{b_0 + b_1(t+1)}$ 。在时段 $t+1$ 和时段 t 的值之间的比率为 $y_{t+1}/y_t = e^{b_0 + b_1(t+1)} / e^{b_0 + b_1(t)} = e^{b_1}$ 。因此，时间序列在两个相邻时间段增长的比率总是相同的： $(y_{t+1} - y_t) / y_t = y_{t+1}/y_t - 1 = e^{b_1} - 1$ 。[⊖]因此，指数增长是以恒定速率增长的。连续复利是出于数学上方便的一种处理方式，该方式使得我们能够将等式重新写为一种容易估计的形式。

如果我们在式 (10-2) 两边同时取自然对数，那么将得到下面的等式：

$$\ln y_t = b_0 + b_1 t, t = 1, 2, \dots, T$$

因此，如果一个时间序列以指数速率增长，那么我们可以对于那个序列的自然对数序列用线性趋势进行建模。[⊙]当然，没有一个时间序列是恰好以一个指数速率增长的。因此，如果我们想要使用一个对数 - 线性模型 (log-linear model)，我们必须估计下面的等式：

$$\ln y_t = b_0 + b_1 t + \epsilon_t, t = 1, 2, \dots, T \tag{10-3}$$

注意到这个等式是关于系数 b_0 和 b_1 线性的。与线性趋势模型中 y_t 的趋势预测值为 $\hat{b}_0 + \hat{b}_1 t$ 不同，对数 - 线性趋势模型中 y_t 的趋势预测值为 $y_t = e^{\hat{b}_0 + \hat{b}_1 t}$ ，因为 $e^{\ln y_t} = y_t$ 。

再观察一下式 (10-3)，我们看到对数 - 线性模型预测 $\ln y_t$ 将以 b_1 的速率从一个时间段增长到下一个时间段。该模型预测 y_t 的常数增长率为 $e^{b_1} - 1$ 。例如，如果 $b_1 = 0.05$ ，那么每期 y_t 的预测增长率为 $e^{0.05} - 1 = 0.051271$ 或者 5.13%。相反，线性趋势模型 (式 (10-1)) 预测 y_t 是以一个固定的数值从这一期增长到下一期的。

例 10-2 展现了线性趋势模型中非随机残差的问题，而例 10-3 则用对数 - 线性模型的设定对于例 10-2 中的数据进行了拟合。

例 10-2 英特尔 (Intel) 公司季度销售额的线性趋势回归

在 2000 年 1 月，技术分析师雷·本尼迪克特 (Ray Benedict) 想要利用式 (10-1) 来拟合英特尔公司季度销售额的数据，如图 10-5 所示。他使用了从 1985 年第一季度到 1999 年第四季度的 60 个英特尔公司销售额的观测值来对于线性趋势回归模型 $y_t = b_0 + b_1 t + \epsilon_t$ 进行估计， $t = 1, 2, \dots, 60$ 。表 10-2 给出了该等式的估计结果。

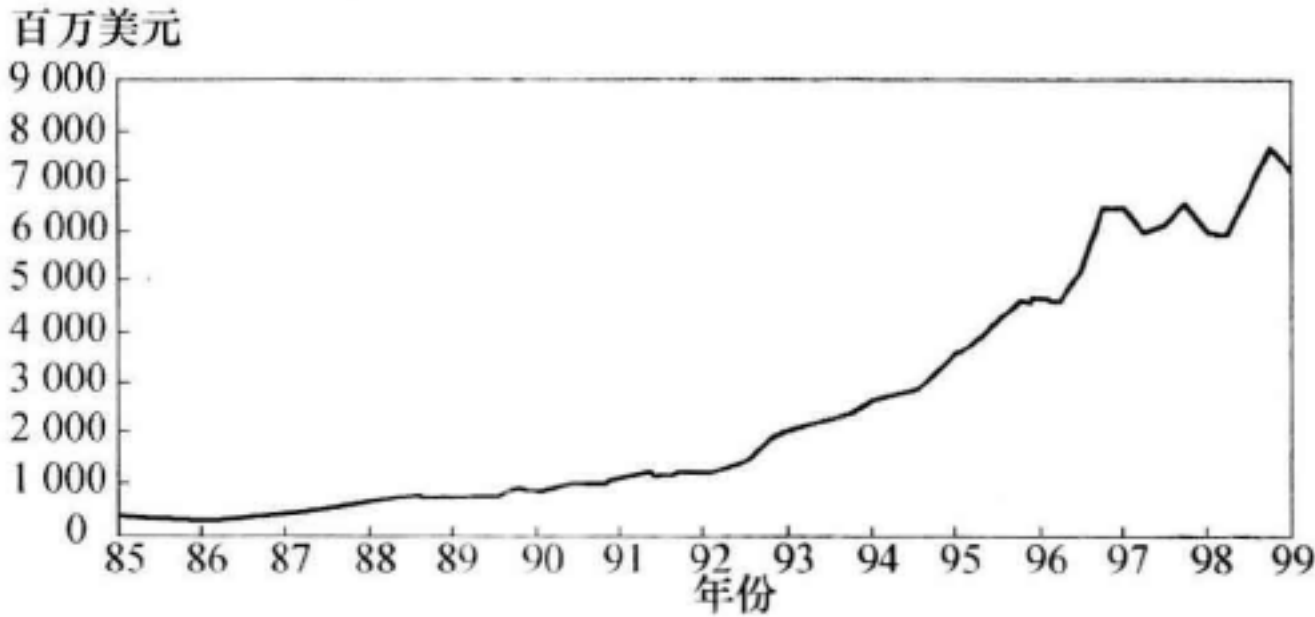


图 10-5 英特尔公司季度销售额

资料来源：Compustat.

⊖ 例如，如果我们使用年作为时间单位，对于一个特定序列的 $e^{b_1} = 1.04$ ，那么该序列就会以每年 $1.04 - 1 = 0.04$ 或者 4% 的速率增长。
⊙ 一个指数增长率是指以连续复利计算的复合增长率。我们在货币的时间价值一章中对于连续复利进行了讨论。

乍看之下，表 10-2 中所给出的结果看上去非常合理：截距和趋势系数都在统计上非常显著。然而，当本尼迪克特将英特尔公司销售额的数据和趋势线画在一张图上时，他却看到另外一幅景象。如图 10-6 所示，在 1989 年之前趋势线持续地低于销售额。在 1989 ~ 1996 年之间，该趋势线持续地高于销售额，但是在 1996 年之后，趋势线又一次持续地低于销售额。

表 10-2 英特尔公司销售额线性趋势模型的估计结果

回归统计量			
R 平方	0.877 4		
标准误	871.6858		
观测数	60		
杜宾 - 沃特森	0.13		
	系数	标准误	t 统计量
截距	-1 318.772 9	227.958 5	-5.785 2
趋势	132.400 5	6.499 4	20.3712

资料来源：Compustat.

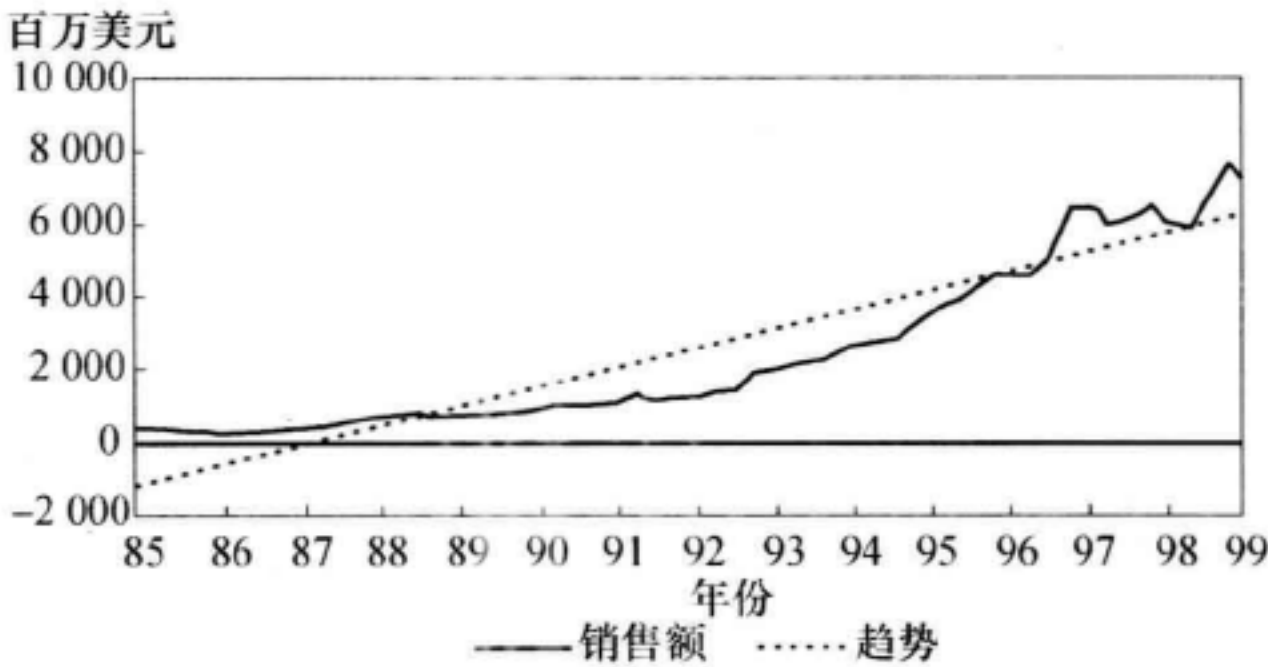


图 10-6 带有趋势的英特尔公司季度销售额

资料来源：Compustat.

回忆一下，一个关于回归模型的重要假设：回归误差在观测值之间应该是不相关的。然而，如果一个趋势持续地高于或者低于时间序列值，那么残差（时间序列和趋势之间的差别）就会是序列相关的。图 10-7 给出了用原始数据估计的线性趋势模型的残差（销售额和趋势之间的差别），该图显示出残差是持续相关的。因为趋势模型误差的持续序列相关，所以即使该等式的 R^2 很高 (0.88)，使用线性趋势来拟合英特尔公司销售额也将是不恰当的。在这里，不相关残差误差的假设被违背了。因为因变量和自变量不像横截面回归中是两个不相同的变量，所以该假设的违背所造成的后果将会是非常严重的，而这使得我们要去寻找更好的模型。

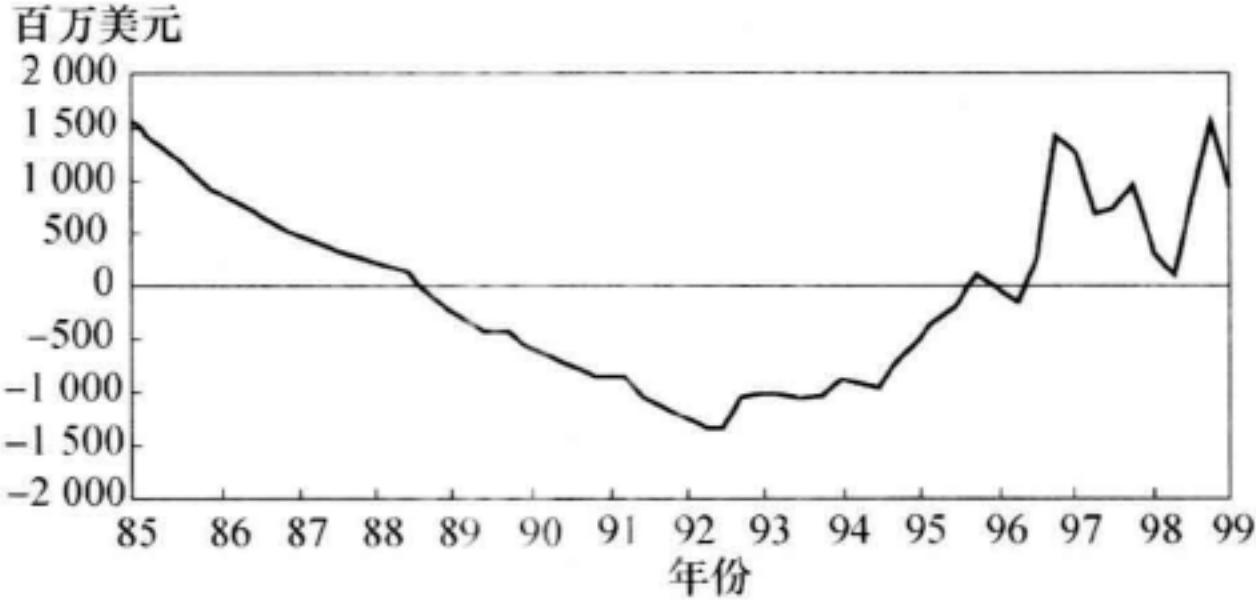


图 10-7 预测带有趋势的英特尔销售额的残差

资料来源：Compustat.

例 10-3 英特尔公司季度销售额的对数 - 线性回归

在例 10-2 中我们已经拒绝了线性趋势模型，现在技术分析师本尼迪克特将对从 1985 年第一季度到 1999 年第四季度的英特尔公司季度销售额尝试另一个不同的模型。在图 10-5 中所显示出的数据点的弯曲形状暗示我们，可能用一个指数曲线可以对于数据进行拟合。于是，他估计了如下的线性方程：

$$\ln y_t = b_0 + b_1 t + \epsilon_t, t = 1, 2, \dots, 60$$

该等式似乎要比式 (10-1) 对于销售额数据拟合得更好。如表 10-3 所示，该等式的 R^2 为 0.98（式 (10-1) 的 R^2 为 0.88）。0.98 的 R^2 意味着英特尔公司销售额自然对数变动的 98% 可以被线性趋势所解释。

图 10-8 显示了用线性趋势拟合英特尔公司销售额自然对数的好坏情况。在样本区间中，销售额数据的自然对数非常接近于线性趋势，并且对数销售额没有长期地高于或者低于趋势。因此，对数 - 线性趋势模型看上去比线性趋势模型更适合来对于英特尔公司销售额进行建模。

本尼迪克特应该如何利用式 (10-3) 的估计结果来预测未来英特尔公司的销售额呢？假设本尼迪克特想要预测 2000 年第一季度英特尔公司的销售额 ($t=61$)。估计值 $\hat{b}_0$ 为 5.552 9，而估计值 $\hat{b}_1$ 为 0.060 9。因此，所估计的模型预测 $\ln \hat{y}_{61} = 5.552 9 + 0.060 9 (61) = 9.267 3$ ，而预测的销售额为 $\hat{y}_{61} = e^{\ln \hat{y}_{61}} = e^{9.267 3} = \$10\,585.63$ （百万）。[⊖]

表 10-3 用线性趋势模型对对数正态英特尔公司销售额的估计

回归统计量			
	R 平方	0.983 1	
	标准误	0.140 7	
	观测数	60	
	杜宾 - 沃特森	0.30	
	系数	标准误	t 统计量
截距	5.552 9	0.036 8	150.980 9
趋势	0.060 9	0.001 0	58.068 0

资料来源：Compustat.

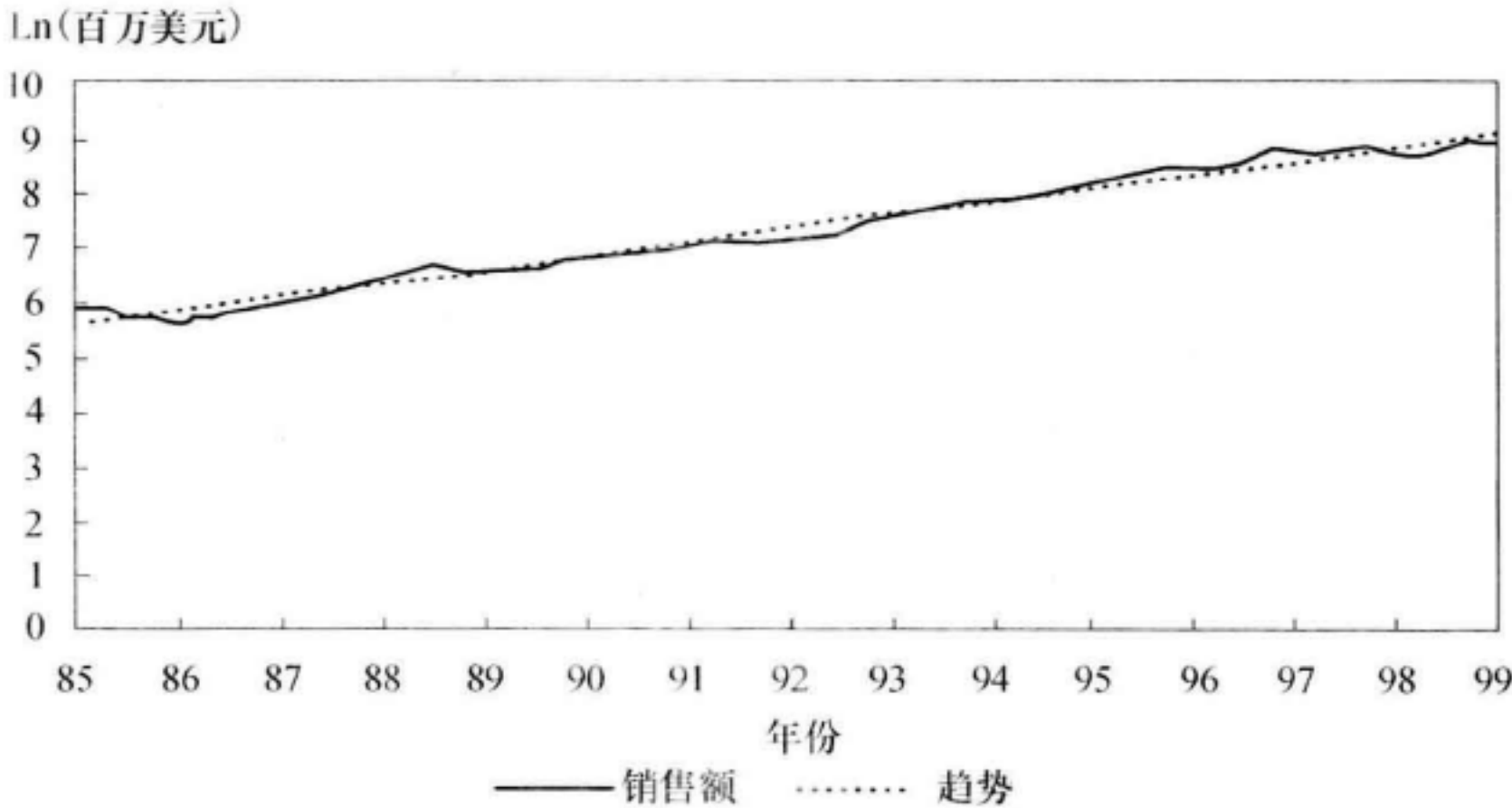


图 10-8 英特尔公司季度销售额的自然对数

⊖ 注意到 $\hat{b}_1 = 0.060 9$ 表明英特尔公司销售额每季度的指数增长率将为 6.28% ($e^{0.060 9} - 1 = 0.062 793$)。

该预测与线性趋势模型的预测有什么不同呢？表 10-2 给出了线性趋势模型的结果，其估计值是 $\hat{b}_0 = -1318.7729$ ，而估计值 $\hat{b}_1$ 是 132.4005。因此，如果我们利用线性趋势模型来预测 2000 年第一季度 ($t = 61$) 英特尔公司的销售额，该预测值为 $\hat{y}_{61} = -1318.7729 + 132.4005(61) = \6757.66 (百万)。这个预测值远低于由对数 - 线性回归模型所得到的结果。在本章之后的部分，我们将考察我们是否能建立一个比仅仅使用对数线性趋势模型更好的模型来对于英特尔公司季度销售额建模。

10.3.3 趋势模型和误差项相关性检验

线性趋势模型和对数 - 线性趋势模型都是单变量回归模型。如果它们正确设定的话，那么回归模型的假设一定会得到满足。特别地，单期回归误差一定和所有其他期的回归误差无关。[Ⓐ]在前一节的例 10-2 中，我们可以从对于残差图像 (图 10-7) 的直观检验中推断出这一假设是否被明显地违背了。例 10-3 的对数 - 线性趋势模型似乎能更好地拟合数据，但是我们仍需要确定不相关误差的假设是否得到满足。为了正规地从理论上回答这个问题，我们必须对残差进行杜宾 - 沃特森检验。

在回归分析一章中，我们介绍了如何应用杜宾 - 沃特森统计量来检验回归误差是否是序列相关的。比如，如果例 10-1 和例 10-3 所给出的趋势模型确实刻画了通货膨胀和英特尔公司销售额时间序列的行为，那么这两个模型的杜宾 - 沃特森统计量就不应该显著地异于 2.0。不然，模型中的误差要么是正序列相关，要么是负序列相关。

在例 10-1 中，对月度 CPI 通货膨胀率的线性趋势进行估计后，我们得到杜宾 - 沃特森统计量为 1.38。这一结果显著不同于 2.0 吗？为了回答这个问题，我们需要检验不存在正序列相关的原假设。对于一个有 192 个观测值的样本和 1 个自变量的模型来说，杜宾 - 沃特森检验统计量在 0.05 显著性水平下的临界值 d_l 要高于 1.65。因为杜宾 - 沃特森统计量的值 (1.38) 要低于这个临界值，所以我们拒绝误差中不存在正相关性的原假设。我们认为使用线性趋势模型来对通货膨胀建模的回归等式，在其误差中存在正的序列相关性。[Ⓑ]我们将需要一个不同类型的回归模型来对此进行刻画，因为原来的模型违背了误差中不存在序列相关性的最小二乘的假设。

在例 10-3 中，对于英特尔公司销售额的自然对数采用线性趋势模型估计所得到的杜宾 - 沃特森统计量为 0.30。假设我们想要检验不存在正序列相关的原假设。在 0.05 显著性水平下，临界值 d_l 为 1.55。杜宾 - 沃特森统计量 (0.30) 的值低于这个临界值，所以我们可以拒绝误差中不存在正序列相关的原假设。我们认为对于英特尔公司季度销售额的自然对数用线性趋势建模的回归等式存在正序列相关的误差。因此，对于这个序列来说，我们也需要建立一个不同类型的模型。

总的来说，我们认为趋势模型有时会具有误差呈现序列相关性的局限。该序列相关性的存在意味着我们要对于这些时间序列建立比趋势模型更好的预测模型。

Ⓐ 注意到时间序列的观测值，不同于横截面的观测值，它是具有逻辑上的先后顺序的：这些观测值一定是按照相关时间段的先后顺序依次出现的。例如，我们不应该利用顺序被打乱了的 CPI 序列的观测值来对于通货膨胀率进行预测，因为诸如自变量的增长这些与时间有关的变动模式，会对所估计的回归系数的统计性质产生不利影响。

Ⓑ 杜宾 - 沃特森统计量在统计上显著的较小值意味着存在正的序列相关性；而显著的较大值表明存在负的序列相关性。这里，DW 统计量 1.38 意味着存在正的序列相关性。更多相关内容参见回归分析一章。

10.4 自回归时间序列模型

对数-线性模型刻画了时间序列的一个重要特征,这个特征同时也是一般时间序列的一个重要特征,就是当期的值和上一期的值有关。例如,英特尔公司当期的销售额与其上一期的销售额有关。自回归模型 (Autoregressive Model, AR 模型),一个用时间序列的历史值来对于时间序列进行解释的模型,很好地刻画了这一相关关系。当我们使用这类模型时,我们可以抛弃传统的将 y 作为因变量,而将 x 作为自变量的记法,因为我们不再能做出这样的区分。这里,我们将简单地统一使用 x_t 来表示回归中的变量。例如,式 (10-4) 给出了关于变量 x_t 的一阶自回归模型, AR(1):

$$x_t = b_0 + b_1 x_{t-1} + \epsilon_t \quad (10-4)$$

因此,在 AR(1) 模型中,我们只使用 x_t 最近一期的历史值来预测当前 x_t 的值。一般地,一个关于变量 x_t 的 p 阶自回归模型, AR(p), 可以表示为

$$x_t = b_0 + b_1 x_{t-1} + b_2 x_{t-2} + \cdots + b_p x_{t-p} + \epsilon_t \quad (10-5)$$

在这个等式中, p 个 x_t 的历史值被用来预测 x_t 的当前值。在下一节中,我们将讨论包含作为自变量的因变量滞后值的时间序列模型的一个重要假设。

10.4.1 协方差平稳序列

注意到式 (10-4) 中的自变量 (x_{t-1}) 是一个随机变量。这个事实可能看上去是一个细微的数学设定,但是事实上并非如此。如果我们使用普通最小二乘法来估计式 (10-4), 当我们具有一个随机分布的自变量,而该自变量是因变量的滞后值,那么我们的统计推断可能是无效的。为了使我们能做出有效的统计推断,我们必须做出一个时间序列分析中的重要假设:我们必须假设我们所要建模的时间序列是协方差平稳的 (covariance stationary)。^①

一个时间序列是协方差平稳的意味着什么呢? 这个基本概念为: 当一个时间序列的数字特征,诸如均值和方差,不是随时间变化的,那么该时间序列就是协方差平稳的。一个协方差平稳的时间序列必须满足三个重要条件。^② 首先,时间序列的期望值必须在所有时间段是一个有限的常数: $E(y_t) = \mu$ 和 $|\mu| < \infty, t = 1, 2, \cdots, T$ 。其次,时间序列的方差必须在所有时间段也是一个有限的常数。最后,对于固定期数之前或之后的时间序列和其自身的协方差必须在所有时间段是一个有限的常数。第二和第三个条件可以统一表示为如下形式:^③

$$\text{Cov}(y_t, y_{t-s}) = \lambda, |\lambda| < \infty, t = 1, 2, \cdots, T; s = 0, \pm 1, \pm 2, \cdots \pm T$$

式中, λ 表示一个常数。如果一个时间序列不是协方差平稳的,但是我们仍旧用式 (10-4) 对其建模,那么会发生什么呢? 这时,估计的结果将不会具有任何经济学含义。对于一个非平稳的时间序列,用式 (10-4) 的回归进行估计将会得到虚假的结果。特别地, b_1 的估计值将会是有偏

① “弱平稳”是协方差平稳的一个同义词。注意,术语“平稳的”或者“平稳性”通常都指的是“协方差平稳”。你也可能遇到约束更强的概念,即“严格”平稳,但是,这个概念的实际应用较少。更多内容,参见 Diebold (2004)。

② 在第一个条件中,我们将使用绝对值来排除均值为负且没有下界(负无穷大)的情况。

③ 当该等式中的 s 等于 0 的时候,它施加了这样一个条件,即时间序列的方差是有限的。得到这一条件的原因是一个随机变量和自身的协方差就是其方差: $\text{Cov}(y_t, y_t) = \text{Var}(y_t)$ 。

的，并且任何假设检验将是无效的。

我们如何判断一个时间序列是否是协方差平稳的呢？我们通常可以通过观察时间序列的图像来回答这个问题。如果图像显示出在各个时间段上大致相同的均值和方差，并且没有显著的季节性效应，那么我们可以假设该时间序列是协方差平稳的。

一些时间序列，如我们在图 10-1 和图 10-4 所看到的，是协方差平稳的。而如在图 10-3 所给出的通货膨胀率数据似乎也大致在样本时间段上具有相同的均值和方差。然而，我们在商业和投资中所遇到的许多时间序列一般都不是协方差平稳的。例如，许多时间序列表现出稳定的增长（或下降），因此，它们的均值不是常数，这意味着它们是不平稳的。作为一个例子，图 10-5 中英特尔公司季度销售额的时间序列明显地表现出随着时间递增的均值。因此，英特尔公司季度销售额不是协方差平稳的。[⊖]宏观经济时间序列诸如那些与收入和消费相关的序列通常也具有很强的趋势性。具有季节性的时间序列（每年以规则的模式变动）也具有不是常数的均值，正如我们之后要讨论的其他类型的时间序列一样，它也不是协方差平稳的。[⊗]

图 10-2 显示月度零售额（没有经过季节调整）同样也不是协方差平稳的。在 12 月的销售额总是远高于其他月份的销售额（这些值是数据中通常的高峰），而 1 月的销售额总是远低于其他的月份（这些值通常是在 12 月高峰之后的低谷）。总的来说，销售额总是随时间而增长的，所以销售额的均值不是恒定的。

在本章后面的部分，我们将说明，我们通常可以将一个非平稳的时间序列转化为一个平稳的时间序列。但是无论一个平稳的时间序列是与生俱来的还是经过后天转化而来的，我们都要注意这样一点：过去的平稳性并不能保证未来的平稳性。总是存在这样的可能性，即设定良好的模型会在世界的实际状况发生变化的情况下失效，而在这时，我们就需要去寻找另一个不同的模型来刻画新的时间序列。

10.4.2 检测自回归模型中的序列相关误差

如果时间序列是协方差平稳的，而且误差是不相关的，那么我们就可以利用普通最小二乘法来对自回归模型进行估计。不幸的是，当自变量包含了自变量的历史值时，我们之前对于序列相关性的检验（杜宾 - 沃特森统计量）就会是无效的。因此，对于大部分时间序列模型，我们无法使用杜宾 - 沃特森统计量。而幸运的是，我们可以使用其他检验来确定一个时间序列模型中的误差是否是序列相关的。在这些检验中的一个揭示了误差项自相关性是否显著地异于 0。这个检验是关于残差自相关系数和残差自相关系数标准误的 t 检验。作为该检验的背景知识，我们下面先讨论一般的自相关系数的概念，然后再转向对于残差自相关系数的介绍。

一个时间序列的自相关系数（autocorrelations）是指时间序列中的随机变量与其自身的历史值之间的相关系数。相关系数的阶数给定为 k ，其中 k 表示滞后期数。当 $k=1$ 时，自相关系数给出的是某一期变量与其之前一期变量间的相关系数。例如， k 阶自相关系数（ k th order autocorrelation）（ ρ_k ）为

$$\rho_k = \frac{\text{Cov}(x_t, x_{t-k})}{\sigma_x^2} = \frac{E(x_t - \mu)(x_{t-k} - \mu)}{\sigma_x^2}$$

⊖ 一般地，任一可以用线性或者对数 - 线性趋势模型进行准确描述的时间序列都不是协方差平稳的，虽然对于其原始序列进行某种变化之后的序列可能是协方差平稳的。
⊗ 特别地，随机游走不是协方差平稳的。

注意, 我们有关系式 $\text{Cov}(x_t, x_{t-k}) \leq \text{Var}(x_t)$, 等号当且仅当 $k=0$ 时成立。这意味着 ρ_k 的绝对值小于或者等于 1。

当然, 我们不能直接观测到自相关系数 ρ_k 。所以, 我们必须估计它们。因此, 我们将 x_t 的期望值 μ 用它的估计值 $\bar{x}$ 来代替, 并计算相关系数的估计值。于是, 时间序列 x_t 的 k 阶自相关系数的估计值 (我们将其记为 $\hat{\rho}_k$) 为

$$\hat{\rho}_k = \frac{\sum_{t=k+1}^T [(x_t - \bar{x})(x_{t-k} - \bar{x})]}{\sum_{t=1}^T (x_t - \bar{x})^2}$$

类比于时间序列自相关系数的定义, 我们可以将时间序列模型残差项的自相关系数定义为^①

$$\begin{aligned}\rho_{\epsilon,k} &= \frac{\text{Cov}(\epsilon_t, \epsilon_{t-k})}{\sigma_\epsilon^2} \\ &= \frac{E(\epsilon_t - 0)(\epsilon_{t-k} - 0)}{\sigma_\epsilon^2} \\ &= \frac{E(\epsilon_t \epsilon_{t-k})}{\sigma_\epsilon^2}\end{aligned}$$

式中, E 表示期望值。我们假设时间序列模型误差项的期望值等于 0。^②

我们能通过检验误差项的自相关系数 (误差项自相关系数 (error autocorrelations)) 是否显著地不同于 0 来确定我们所使用的时间序列模型是否正确。如果这些自相关系数确实显著地不同于 0, 那么我们所使用的模型就是设定错误的。我们利用残差的样本自相关系数 (残差自相关系数 (residual autocorrelations)) 和它们的样本方差来估计误差的自相关系数。

关于特定滞后期的误差项自相关系数等于 0 的原假设的检验基于的是同样滞后期的残差自相关系数和残差自相关系数的标准误, 其等于 $1/\sqrt{T}$, 其中 T 是时间序列的观测数。^③ 因此, 如果我们的时间序列具有 100 个观测值, 那么每个自相关系数估计值的标准误就是 0.1。我们可以通过将给定滞后期的残差自相关系数除以其标准误 ($1/\sqrt{T}$) 来计算得到关于该滞后期的误差相关系数等于 0 的原假设的 t 检验量。

那么, 我们如何利用误差项自相关系数所提供的信息来确定自回归时间序列模型是否是正确设定的呢? 我们可以使用简单的三步法。第一步, 估计一个特定的自回归模型, 比如说 AR(1) 模型。第二步, 计算该模型残差的自相关系数。^④ 第三步, 通过检验来观察残差的自相关系数是否显著地不同于 0。如果显著性检验表明残差的自相关系数显著地不同于 0, 那么模型就是设定错误的; 我们可能要采用我们不久之后将要介绍的方法来对模型进行修正。^⑤ 我们现在给出一个例子

① 当我们没有任何修饰地提到自相关系数时, 我们指的是时间序列本身的自相关系数, 而不是指误差项或者残差的自相关系数。

② 这一假设类似于之前两章中我们对于误差项期望值所做出的假设。

③ 这一计算是在 Diebold (2004) 中推导出来的。

④ 我们可以用大部分统计软件包非常方便地计算出这些残差的自相关系数。例如, 在微软的 Excel 中, 为了计算一阶的残差自相关系数, 我们可以计算由第 1 个到第 $T-1$ 个观测值所得到的残差和由第 2 个到第 T 个观测值所得到的残差之间的相关系数。

⑤ 通常, 经济计量学家使用另外的方法来对于残差自相关系数的显著性进行检验。例如, Box-Pierce Q 统计量常被用来检验所有残差的自相关系数等于 0 的联合假设。更多相关讨论参见 Diebold (2004)。

来展现这个三步法是如何进行的。

例 10-4 预测英特尔公司的毛利润率

在考察了英特尔公司销售额时间序列的模型之后, 分析师雷·本尼迪克特决定使用一个时间序列模型来预测英特尔公司的毛利润率 [(销售收入 - 销售成本) / 销售收入]。他对于因变量的观测时间段是从 1985 年第二季度到 1999 年第四季度。他并不知道刻画毛利润率最好的模型是什么, 但是他认为当期值与前一期的值有关。他决定从一阶自回归模型, AR(1) 开始着手: $\text{毛利润率}_t = b_0 + b_1 (\text{毛利润率}_{t-1}) + \epsilon_t$ 。表 10-4 给出了该 AR(1) 模型的估计结果, 同时也给出了这个模型残差的自相关系数。

在表 10-4 中我们首先注意到的是回归等式中截距 ($\hat{b}_0 = 0.0834$) 和毛利润率一阶滞后项的系数 ($\hat{b}_1 = 0.8665$) 都是高度显著的。^① 截距的 t 统计量为 2.3, 而毛利润率一阶滞后项的 t 统计量超过了 14。在 59 个观测值和 2 个参数的情况下, 该模型具有 57 个自由度。在 0.05 的显著性水平下, t 统计量的临界值约为 2.0。因此, 本尼迪克特将必须拒绝截距等于 0 ($b_0 = 0$) 和一阶滞后项系数等于 0 ($b_1 = 0$) 的原假设, 而接受这两个系数都不等于 0 的备择假设。但是这些统计量是有效的吗? 当我们检验模型的残差是否是序列相关之后, 我们将了解到这一点。

表 10-4 的底部给出了前四个残差的自相关系数以及这些自相关系数所对应的标准误和 t 统计量。^② 该样本具有 59 个观测值, 所以每个自相关系数的标准误为 $1/\sqrt{59} = 0.1302$ 。表 10-4 显示前四个自相关系数中没有一个是 t 统计量的绝对值大于 1.50。因此, 本尼迪克特可以认为这些自相关系数中没有一个是显著异于 0。结果, 他可以假设残差不是序列相关的, 而且模型的设定是正确的, 而他也可以有效地利用普通最小二乘法来估计自相关模型中的参数和参数的标准误。^③

现在本尼迪克特已经得出该模型的设定是正确的, 那么他如何利用该模型来对于英特尔公司下一期的毛利润率进行预测呢? 所估计的等式为 $\text{毛利润率}_t = 0.0834 + 0.8665 (\text{毛利润率}_{t-1}) + \epsilon_t$ 。在任何时间段, 误差项的期望值都为 0。因此, 该模型预测 $t+1$ 期的毛利润率将为 $\text{毛利润率}_{t+1} = 0.0834 + 0.8665 (\text{毛利润率}_t)$ 。例如, 如果本季度毛利润率为 55% (0.55), 那么该模型预测下一季度的毛利润率将增长到 $0.0834 + 0.8665 (0.55) = 0.5600$ 或者 56%。另一方面, 如果当前的毛利润率为 65% (0.65), 那么模型预测下一期的毛利润率将下降到 $0.0834 + 0.8665 (0.65) = 0.6467$ 或者 64.67%。正如我们在下一节中所示, 如果毛利润率低于一定水平 (62.50), 那么该模型预测的毛利润率将会上升, 而如果毛利润率高于该水平, 那么该模型预测的毛利润率将会下降。

① 时间序列的一阶滞后项是上一期时间序列的值。

② 对于没有经过季节调整的数据, 分析师通常计算与一年中的观测数目相同多的自相关系数 (例如, 对于季度数据, 我们就计算 4 个自相关系数)。所计算的自相关系数的数目通常也会依赖于样本的大小, 正如 Diebold (2004) 所讨论的。

③ 统计学家拥有许多其他检验时间序列模型中残差序列相关性的方法。更多具体内容参见 Diebold (2004)。

表 10-4 自回归模型：AR(1) 模型，英特尔公司季度毛利润率的观测值（1985 年 4 月 ~ 1999 年 12 月）

回归统计量			
	R 平方	0.778 4	
	标准误	0.0402	
	观测数	59	
	杜宾 - 沃特森	1.8446	
	系数	标准误	t 统计量
截距	0.083 4	0.036 7	2.270 5
滞后 1 期	0.866 5	0.061 2	14.149 3
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	0.067 7	0.130 2	0.519 7
2	-0.192 9	0.130 2	-1.4814
3	0.054 1	0.130 2	0.415 2
4	-0.149 8	0.130 2	-1.150 7

资料来源：Compustat.

10.4.3 均值回复

我们认为如果当一个时间序列高于其均值时，其值有下降的倾向，而当它低于其均值，该时间序列又有上升的倾向时，该时间序列就表现出均值回复（mean reversion）的特征。非常类似于受恒温器所调节的室内温度，一个均值回复的时间序列倾向于回归到其长期的均值。那么我们如何确定该时间序列倾向于回归到数值呢？如果一个时间序列当前在其均值回复的水平上，那么模型将会预测下一期该时间序列的值将保持不变。在其均值回复的水平上，我们具有关系式 $x_{t+1} = x_t$ 。对于 AR(1) 模型 ($x_{t+1} = b_0 + b_1x_t$) 来说，等式 $x_{t+1} = x_t$ 意味着水平 $x_t = b_0 / (1 - b_1)$ ，或者该均值回复水平 x_t 可以表示为

$$x_t = \frac{b_0}{1 - b_1}$$

因此，如果时间序列的当前值为 $b_0 / (1 - b_1)$ 的话，那么，AR(1) 模型预测该时间序列将保持不变。如果其当前值低于 $b_0 / (1 - b_1)$ ，那么 AR(1) 模型预测该时间序列将上升，而如果其当前值高于 $b_0 / (1 - b_1)$ ，那么 AR(1) 模型预测该时间序列将下降。

在英特尔公司毛利润率的例子中，由表 10-4 所给出的模型的均值回复水平为 $0.083\ 4 / (1 - 0.866\ 5) = 0.6250$ 。如果当前的毛利润率高于 0.625 0，那么模型预测毛利润率将会在下一期下降。如果当前的毛利润率低于 0.6250，那么模型预测毛利润率将会在下一期上升。正如我们将在之后讨论的，所有的协方差平稳的时间序列都具有一个有限的均值回复水平。

10.4.4 多期预测和预测的链式法则

通常，金融分析师们想要对于多期进行预测。例如，我们可能想要利用季度销售额模型预测接下来四个季度公司的销售额。为了使用时间序列模型来对于多期进行预测，我们必须考察如何利用 AR(1) 模型进行多期预测。对于 AR(1) 模型中 x_t 的向前一期的预测可以表达如下：

$$\hat{x}_{t+1} = \hat{b}_0 + \hat{b}_1x_t \tag{10-6}$$

如果我们想要利用 AR(1) 模型预测 x_{t+2} ，那么我们的预测将基于

$$\hat{x}_{t+2} = \hat{b}_0 + \hat{b}_1 x_{t+1} \quad (10-7)$$

不幸的是，我们在 t 时刻不知道 x_{t+1} 的值，所以我们无法直接利用式 (10-7) 来做出向前两期的预测。然而，我们可以利用我们对于 x_{t+1} 的预测值以及 AR(1) 模型来对于 x_{t+2} 进行预测。预测的链式法则 (chain rule of forecasting) 是指一个预测过程，在这个过程中，由预测方程所得到的下一期的预测值将被再次代入预测方程之中以获得两期之后的预测值。利用预测的链式法则，我们可以将 x_{t+1} 的预测值代入式 (10-7) 中，从而得到 $\hat{x}_{t+2} = \hat{b}_0 + \hat{b}_1 \hat{x}_{t+1}$ 。我们已经从式 (10-6) 的向前一期预测中知道了 x_{t+1} 。所以现在我们有了一个简便的预测 x_{t+2} 的方法。

多期预测比单期预测具有更大的不确定性，因为每一个预测期都有不确定性。例如，在预测 x_{t+2} 时，我们首先会产生与利用 x_t 来预测 x_{t+1} 相关的不确定性，其次，我们又会产生与利用 x_{t+1} 的预测值来预测 x_{t+2} 相关的不确定性。一般地，一个预测值距当前的时刻越远，其不确定性也就越大。[⊖]

例 10-5 英特尔公司毛利润率的多期预测

假设在 2000 年年初，我们想要利用表 10-4 给出的模型来预测两期的英特尔公司的毛利润率。假设当期英特尔公司毛利润率为 65%。由模型所得到的向前一期的英特尔公司毛利润率预测值为 $0.6467 = 0.0834 + 0.8665 (0.65)$ 。通过将向前一期的预测值 0.6467 代回到回归等式中，我们就可以得到下面的向前两期的预测： $0.6438 = 0.0834 + 0.8665 (0.6467)$ 。因此，如果当期的英特尔公司毛利润率为 65%，那么两期后该模型预测英特尔公司的毛利润率将为 64.38%。

例 10-6 对美国 CPI 通货膨胀率建模

分析师莉塞特·米勒被要求对于美国月度通货膨胀率进行时间序列建模。通货膨胀率对于通货膨胀率的预期都会对债券收益率产生显著的效应。莉塞特·米勒选取了自 1971 年 1 月开始的年化的月度 CPI 百分比变动的数据。那么，米勒应该使用哪一个模型呢？

选择模型的过程和对于英特尔公司毛利润率建模的例 10-4 的过程是一样的。米勒所估计的第一个模型是 AR(1) 模型，她将前一个月的通货膨胀率作为自变量：通货膨胀率 _{t} = $b_0 + b_1$ 通货膨胀率 _{$t-1$} + ϵ_t , $t = 1, 2, \dots, 359$ 。为了估计这个模型，她使用了从 1971 年 1 月到 2000 年 12 月的月度 CPI 通货膨胀率数据（用 $t = 1$ 表示 1971 年 2 月）。表 10-5 给出了这个模型的估计结果。

如表 10-5 所示，截距 ($\hat{b}_0 = 1.9658$) 和通货膨胀率的一阶滞后值 ($\hat{b}_1 = 0.6175$) 在统计上都是非常显著的，因为其具有很大的 t 统计量。在 359 个观测值和 2 个参数的情况下，该模型的自由度为 357。在 0.05 显著性水平下， t 统计量的临界值约为 1.97。因此，米勒能够拒绝截距等于 0 ($b_0 = 0$) 和一阶滞后项的系数等于 0 的单个原假设，而接受单个系数不等于 0 的备择假设。

⊖ 如果一个预测模型是良好设定的，那么模型的预测误差将不是序列相关的。如果每期预测误差不是序列相关的，那么多期预测的方差将高于单期预测的方差。

这些统计量是有效的吗？当米勒检验这个模型的残差是否是序列相关之后，她将了解到这一点。在这个样本中有 359 个观测值，所以每个自相关系数估计值的标准误为 $1/\sqrt{359} = 0.0528$ 。 t 统计量的临界值为 1.97。因为第一个和第二个自相关系数估计量的 t 统计量绝对值都大于 1.97，所以米勒认为自相关系数显著地不同于 0。因此，这个模型是设定错误的，因为其残差是序列相关的。

表 10-5 年化的月度 CPI 通货膨胀率：AR(1) 模型月度观测值（1971 年 2 月~2000 年 12 月）

回归统计量			
	<i>R</i> 平方	0.3808	
	标准误	3.4239	
	观测数	359	
	杜宾-沃特森	2.3059	
	系数	标准误	<i>t</i> 统计量
截距	1.9658	0.2803	7.0119
滞后 1 期	0.6175	0.0417	14.8185
残差的自相关系数			
滞后期限	自相关系数	标准误	<i>t</i> 统计量
1	-0.1538	0.0528	-2.9142
2	0.1097	0.0528	2.0782
3	0.0657	0.0528	1.2442
4	0.0920	0.0528	1.7434

U. S. Bureau of Labor Statistics.

如果自回归模型中的残差是序列相关的，那么米勒可以通过对将更多因变量的滞后项作为解释变量的自回归模型进行估计来消除残差的序列相关性。表 10-6 使用与表 10-5 的分析中相同的数据，给出了第二个时间序列模型，AR(2) 模型的估计结果。[⊖]在 358 个观测值和 3 个参数的情况下，这个模型具有 355 个自由度。因为自由度几乎和表 10-5 中所给出的估计量相同，所以在 0.05 显著性水平下的 t 统计量的临界值也几乎是相同的（1.97）。如果米勒用两阶滞后值来估计等式，通货膨胀率_{*t*} = $b_0 + b_1$ 通货膨胀率_{*t-1*} + b_2 通货膨胀率_{*t-2*} + ε_t ，那么她将会发现回归模型中的三个系数都显著地不同于 0。表 10-6 的底部显示残差的前四阶自相关系数 t 统计量的绝对值没有一个大于临界值 1.97。因此，米勒无法拒绝单个残差自相关系数等于 0 的原假设。故她认为该模型的设定正确，因为她没有发现残差序列相关的证据。

表 10-6 年化的月度 CPI 通货膨胀率：AR(2) 模型月度观测值（1971 年 3 月~2000 年 12 月）

回归统计量			
	<i>R</i> 倍	0.6479	
	<i>R</i> 平方	0.4197	
	标准误	3.3228	
	观测数	358	
	杜宾-沃特森	2.0582	
	系数	标准误	<i>t</i> 统计量
截距	1.4609	0.2913	5.0147
滞后 1 期	0.4634	0.0514	9.0117
滞后 2 期	0.2515	0.0514	4.8924

⊖ 注意，表 10-6 显示回归中只有 358 个观测值，这是因为额外的通货膨胀率滞后项需要对于表 10-5 的回归中一个月之后的样本进行估计。（在两阶滞后的情况下，1971 年 1 月和 2 月的通货膨胀率必须是已知的，以使得我们能够估计从 1971 年 3 月开始的等式。）

(续)

残差的自相关系数			
滞后期限	自相关系数	标准误	t统计量
1	-0.032 0	0.052 9	-0.604 8
2	-0.098 2	0.052 9	-1.857 4
3	-0.011 4	0.052 9	-0.215 0
4	0.032 0	0.052 9	0.605 3

资料来源：U. S. Bureau of Labor Statistics.

在上一个例子中，分析师选择 AR (2) 模型，因为 AR (1) 模型的残差是序列相关的。假设在一个给定的月份，已知上个月的通货膨胀率为年度 4%，而再上个月的通货膨胀率为年度 3%。那么，表 10-5 中所给出的 AR (1) 模型预测下一个月的通货膨胀率将约为 $1.9658 + 0.6175(4) = 4.44\%$ ，而表 10-6 中所给出的 AR (2) 模型预测下个月的通货膨胀率将约为 $1.4609 + 0.4634(4) + 0.2515(3) = 4.07\%$ 。如果分析师使用不正确的 AR (1) 模型，那么她预测出的通货膨胀率将比使用 AR (2) 模型高出 37 个基点。这一不正确的预测可能对于其公司的投资决策的质量产生不利的影响。

10.4.5 比较预测模型的表现

一种比较两个模型的预测表现的方法是比较两个模型所预测误差的方差。具有更小预测误差方差的模型将是更精确的模型，而且它也具有更小的时间序列回归的标准误。（该标准误通常在时间序列回归的输出中直接进行报告。）

在不同模型预测准确性的比较中，我们必须区分样本内预测误差和样本外预测误差。样本内预测误差 (In-sample forecast errors) 是拟合时间序列模型的残差。例如，当我们用线性趋势对 1971 年 1 月到 2000 年 12 月的原始通货膨胀数据进行估计，样本内预测误差是 1971 年 1 月到 2000 年 12 月的残差。如果我们使用这个模型来预测该区间之外的通货膨胀率，那么，实际通货膨胀率和预测通货膨胀率的差别是样本外预测误差 (Out-of-sample forecast errors)。

例 10-7 美国 CPI 通货膨胀率的样本内预测比较

在例 10-6 中，分析师比较对于月度美国通货膨胀率的 AR (1) 预测模型和对于月度美国通货膨胀率的 AR (2) 模型，并得出 AR (2) 模型更好。表 10-5 给出了关于通货膨胀率的 AR (1) 模型的标准误为 3.423 9，而表 10-6 给出了 AR (2) 模型的标准误为 3.322 8。因此，AR (2) 模型具有比 AR (1) 模型更小的样本内预测误差的方差，而这是与我们认为 AR (2) 模型更好的想法一致的。该模型的标准误是 AR (1) 模型预测误差的 $3.3228 / 3.4239 = 97.05\%$ 。

通常，我们想要比较所估计的样本区间段之外的不同模型的预测准确性。我们想要比较模型的样本外预测误差。样本外预测误差是重要的，因为未来总是在样本外的。虽然职业预测家会区分样本外和样本内的预测表现，但是许多分析师阅读的文献只包括样本内预测的评估。分析师应该意识到样本外的表现对于评价一个预测模型对于现实世界的贡献是非常重要的。

一般地，我们通过模型均方误差平方根 (root mean squared error, RMSE) (其是误差平方平均值的平方根) 之间的比较来判断样本外的预测模型的预测表现。具有最小 RMSE 的模型将被认为是最准确的。下面的例子给出了 RMSE 在比较预测模型中的计算和使用方法。

例 10-8 美国 CPI 通货膨胀率的样本外预测比较

假设我们想要用 2001 年 1 月到 2002 年 12 月的美国通货膨胀率数据来比较用 AR (1) 和 AR (2) 模型估计出的 1971 ~ 2000 年的美国通货膨胀率的预测准确性。

对于 2001 年 1 月到 2002 年 12 月 中的每个月份，表 10-7 中第一列的数字表示在该月份实际年化的通货膨胀率。第二列和第三列表示之前两个月的通货膨胀率。第四列给出了表 10-5 中所示的 AR (1) 模型的样本外误差。第五列给出了 AR (1) 模型的平方误差。第六列给出了表 10-6 所示的 AR (2) 模型的样本外误差。最后一列给出了 AR (2) 模型的平方误差。表 10-7 的底部给出了平均平方误差和 RMSE。根据这些度量指标，AR (1) 模型在 2001 年 1 月到 2002 年 12 月通货膨胀率的样本外预测上要比 AR (2) 模型更准确一些。AR (1) 模型的 RMSE 仅为 AR (2) 模型 RMSE 的 $3.7474/3.8116=98.32\%$ 。因此，即使 AR (2) 模型在样本内预测上更加准确，但是 AR (1) 在样本外预测上要更准确一些。当然，这里仅是用一个小样本来对于样本外预测的表现进行评估。在这里，虽然我们似乎得到了关于选择 AR (1) 还是 AR (2) 模型的相互冲突的结果，但是我们还必须考虑回归系数的稳定性。我们将在下面一节中继续对于这两个模型进行比较。

表 10-7 样本外预测误差比较：2001 年 1 月 ~ 2002 年 12 月美国 CPI 通货膨胀率（已经经过年化的）

日期 (2001 年)	Infl (t)	Infl (t-1)	Infl (t-2)	AR (1) 误差	平方误差	AR (2) 误差	平方误差
1 月	7.855 6	-0.687 1	0.691 8	6.314 1	39.868 1	6.539 2	42.761 1
2 月	4.904 2	7.855 6	-0.687 1	-1.912 2	3.656 4	-0.024 2	0.000 6
3 月	2.764 8	4.904 2	7.855 6	-2.229 1	4.968 9	-2.944 0	8.666 9
4 月	4.872 9	2.764 8	4.904 2	1.199 9	1.439 9	0.897 6	0.805 8
5 月	5.563 8	4.872 9	2.764 8	0.589 2	0.347 2	1.149 7	1.321 7
6 月	2.044 8	5.563 8	4.872 9	-3.356 4	11.265 6	-3.219 6	10.366 0
7 月	-3.319 2	2.044 8	5.563 8	-6.547 5	42.870 3	-7.126 7	50.789 2
8 月	0.000 0	-3.319 2	2.044 8	0.083 7	0.007 0	-0.436 9	0.190 9
9 月	5.544 6	0.000 0	-3.319 2	3.578 8	12.807 8	4.918 3	24.189 9
10 月	-3.964 2	5.544 6	0.000 0	-9.353 6	87.489 1	-7.994 4	63.911 0
11 月	-2.007 2	-3.964 2	5.544 6	-1.525 2	2.326 1	-3.025 2	9.151 9
12 月	-4.633 6	-2.007 2	-3.964 2	-5.360 0	28.729 9	-4.167 5	17.368 5
2002							
1 月	2.750 5	-4.633 6	-2.007 2	3.645 9	13.292 7	3.941 6	15.536 5
2 月	4.847 6	2.750 5	-4.633 6	1.183 4	1.400 5	3.277 2	10.740 4
3 月	6.961 9	4.847 6	2.750 5	2.002 9	4.011 8	2.563 0	6.569 2
4 月	6.921 8	6.961 9	4.847 6	0.657 3	0.432 0	1.015 8	1.031 9
5 月	0.000 0	6.921 8	6.961 9	-6.239 7	38.933 9	-6.419 0	41.203 4
6 月	0.669 5	0.000 0	6.921 8	-1.296 3	1.680 4	-2.531 9	6.410 5
7 月	1.342 3	0.669 5	0.000 0	-1.036 9	1.075 1	-0.428 8	0.183 9
8 月	4.071 9	1.342 3	0.669 5	1.277 3	1.631 5	1.820 7	3.314 8
9 月	2.010 5	4.071 9	1.342 3	-2.469 4	6.098 1	-1.674 7	2.804 7
10 月	2.007 2	2.010 5	4.071 9	-1.200 0	81.440 0	-1.409 2	1.986 0
11 月	0.000 0	2.007 2	2.010 5	-3.205 1	10.272 8	-2.896 5	8.390 0
12 月	-2.615 7	0.000 0	2.007 2	-4.581 4	20.989 3	-4.581 2	20.987 6
				平均	14.043 1	平均	14.528 4
				RMSE	3.747 4	RMSE	3.811 6

资料来源：U. S. Bureau of Labor Statistics.

10.4.6 回归系数的不稳定性

分析师在对一个时间序列建模时所面临的一个重要问题是：应该使用哪个样本时间段的数据。时间序列模型的回归系数的估计值可能会随着用于估计模型的时间段的改变而产生巨大的变化。通常，利用早期的样本时间段所得到的时间序列模型的回归系数估计值可能会与后期的样本时间段所得到的时间序列模型的估计值有较大的不同。类似地，对于相对较短和相对较长的样本时间段的不同模型，所得到的估计值可能也会不同。更进一步，特定时间序列模型的选择也会依赖于样本时间段。例如，AR(1)模型可能在某个特定样本时间段用来对于公司的销售额建模是适合的，而AR(2)模型可能在样本时间段的早期或者后期是适合的（或者在更长或者更短的样本时间段）。因此，样本时间段的选择在金融时间序列建模中是一个重要的决策。

不幸的是，在经济或者金融理论中通常没有明确给出估计时间序列模型应使用较长还是较短样本时间段的理论。然而，如果我们记住只有协方差平稳的时间序列模型才是有效的，那么我们在选择时也许能够得到一些帮助。例如，我们不应该将固定汇率时间段的数据和浮动汇率时间段的数据放在一起。在这两个时间段的汇率不可能具有相同的方差，因为在浮动汇率制度下的汇率通常要比固定汇率制度下的汇率波动要剧烈。同样地，许多美国分析师认为将20世纪60年代作为模型中美国通货膨胀率或者利率数据中的一部分是不适合的，因为美联储在这个时间段采取了不同的政策制度。确定估计时间序列所使用的适当样本的最好方法是观察数据的散点图，以判断时间序列在估计之前是否是平稳的。如果我们了解到政府政策在某个时点发生了变动，那么我们也应该检验一下，该时点前后的时间序列是否保持相同的关系。

在下面的例子中，我们说明了较长时间段和较短时间段是如何影响使用一阶或者二阶时间序列模型的选择决策的。接着我们说明时间序列模型（以及相关的回归系数）的选择是如何影响我们的预测的。最后，我们讨论哪个样本时间段、哪个模型及其相应的预测是适合该例子来进行时间序列分析的。

例 10-9 美国通货膨胀率时间序列模型的非平稳性

在例10-6中，分析师莉塞特·米勒认为美国CPI通货膨胀率应该用AR(2)模型来建模。她的一个同事检验了她的结果，并质疑其用该模型对于1971年之后的美国通货膨胀率的估计结果，因为美联储的政策在20世纪70年代末、80年代初发生了巨大的变化。他认为从1971~2000年的通货膨胀率时间序列具有两种机制(regimes)或者由两种对应的模型所产生：一种从1971~1984年，而另一种从1985年开始。因此，同事建议米勒对于从1985年开始的美国通货膨胀率估计一个新的时间序列模型。根据他的建议，米勒首先使用从1985~2000年的更短的样本时间段的数据对通货膨胀率用AR(1)模型进行了估计。表10-8给出了她AR(1)模型的估计结果。

表10-8的底部显示AR(1)模型残差的前四阶的自相关系数很小。这些自相关系数中没有有一个 t 统计量显著大于临界值1.97。所以，米勒无法拒绝残差是序列不相关的原假设。AR(1)模型对于1985~2000年的样本时间段是设定正确，因此，我们没有必要再对于AR(2)模型进行估计。这一结果与例10-6中利用1971~2000年所得到的估计结果完全不同。在那个例子中，米勒最初拒绝了AR(1)模型，因为其残差表现出序列相关性。当她使用更大的样本时，AR(2)模型在最初的时间段中比AR(1)模型拟合数据更好。

表 10-8 自回归：AR（1）模型年化的月度 CPI 通货膨胀率（1985 年 2 月~2000 年 12 月）

回归统计量			
R 平方		0.154 0	
标准误		2.464 1	
观测数		191	
杜宾—沃特森		1.918 2	
	系数	标准误	t 统计量
截距	2.137 1	0.285 9	7.474 7
滞后 1 期	0.335 9	0.069 0	4.871 6
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	0.028 4	0.072 4	0.392 2
2	-0.090 0	0.072 4	-1.242 6
3	-0.014 1	0.072 4	-0.195 5
4	-0.029 7	0.072 4	-0.410 3

资料来源：U. S. Bureau of Labor Statistics.

那么，样本时间段的选择对于未来通货膨胀率预测的影响有多大呢？假设在一个给定的月份，年化的通货膨胀率为 4%，而在该月之前的通货膨胀率为 3%。表 10-8 给出的 AR（1）模型预测下个月的通货膨胀率将约为 $2.1371 + 0.3359(4) = 3.48(\%)$ 。因此，利用 1985~2000 年的样本所得到的下个月通货膨胀率的预测值为 3.48%。记得根据例 10-6 的分析，AR（2）模型对于 1971~2000 年样本的下月通货膨胀率的预测值为 4.07%。因此，使用更短样本的设定正确的模型会得到一个几乎低于更长样本的设定正确的模型 0.6% 的通货膨胀预测值。这一差别可能会极大地影响到某一投资决策。

哪个模型是正确的？图 10-9 给出了答案。在 20 世纪 70 年代中期到 20 世纪 80 年代初美国月度通货膨胀率要比 1984 年之后数值更高且波动更剧烈，这说明 1971~2000 年的通货膨胀率很可能不是协方差平稳的时间序列。因此，我们有理由相信数据具有不止一种状态，而米勒应该如同上面所示，分开估计 1985~2000 年的通货膨胀模型。正如例子所示，判断力和经验（诸如对于政府政策变动的了解）在确定一个时间序列如何建模的过程中发挥着重要的作用。仅仅依靠时间序列模型的残差自相关系数无法告诉我们哪一个才是我们分析所需要的正确的样本时间段。

年化通货膨胀率（%）

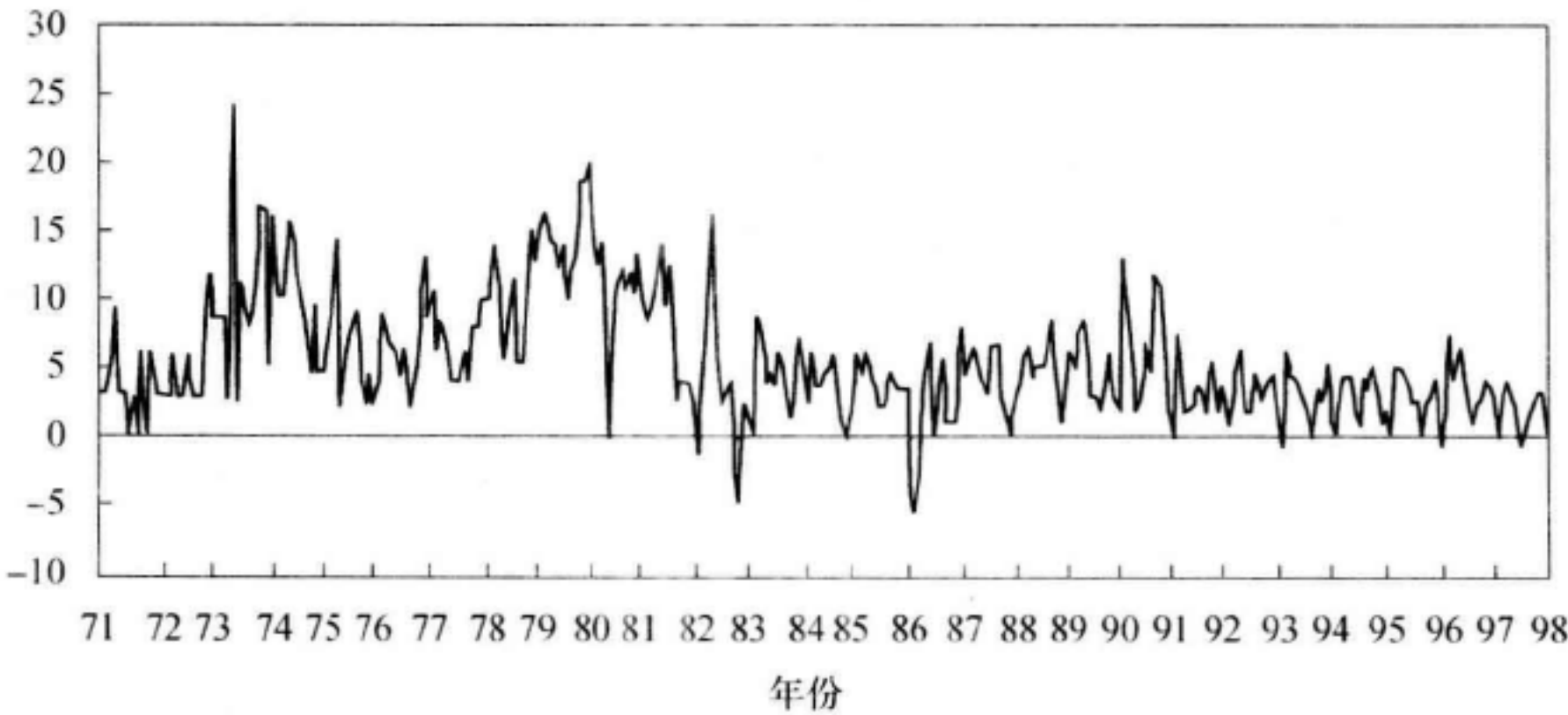


图 10-9 月度 CPI 通货膨胀率

资料来源：U. S. Bureau of Labor Statistics.

10.5 随机游走和单位根

至今,我们已经考察了具有回复到其均值水平倾向的那些时间序列,那些时间序列中的变量在前一期到后一期的变动遵循着一个均值回复的模式。而与此相反的是,现实中还存在着许多金融时间序列,其变动服从一个随机的模式。我们将在下面对于这些“随机游走”进行讨论。

10.5.1 随机游走

随机游走是在金融数据建模中研究最为广泛的一类时间序列模型。一个随机游走 (random walk) 是指这样一种时间序列,其每一期的值都等于上一期的值加上一个不可预测的随机误差。一个随机游走可以用下面的等式来描述:

$$x_t = x_{t-1} + \epsilon_t, E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2, E(\epsilon_t \epsilon_s) = 0 \quad \text{如果 } t \neq s \quad (10-8)$$

式 (10-8) 意味着时间序列 x_t 在每一期等于其前一期加上一个误差项 ϵ_t , 而该误差项具有常数方差, 且与前面各期的误差项不相关。有两个重要的地方需要我们注意。首先, 该等式是 AR (1) 模型 $b_0 = 0, b_1 = 1$ 的特例。[⊖] 其次, ϵ_t 的期望值为 0。因此, 在 $t-1$ 时刻对于 x_t 的最佳预测值为 x_{t-1} 。事实上, 在这个模型中, x_{t-1} 是 $t-1$ 之后每一期 x 的最佳预测值。

随机游走在金融时间序列中非常常见。例如, 许多研究经过检验后发现货币汇率服从一个随机游走过程。与上面第二个观点一致, 一些研究发现复杂的汇率预测模型无法做出超越随机游走模型的预测, 并且未来汇率的最佳预测就是当前的汇率。

不幸的是, 我们无法使用上面所讨论的回归方法来估计一个实际服从随机游走的时间序列的 AR (1) 模型。为了了解为什么是这样的, 我们必须确定为什么一个随机游走不具有有限的均值回复水平或者有限方差。回忆一下, 如果 x_t 是其均值回复水平, 那么 $x_t = b_0 + b_1 x_t$, 或者 $x_t = b_0 / (1 - b_1)$ 。然而, 在一个随机游走中, $b_0 = 0, b_1 = 1$, 所以, $b_0 / (1 - b_1) = 0/0$ 。因此, 一个随机游走具有无法确定的均值回复水平。

一个随机游走的方差是什么呢? 假如在第一期, x_1 的值为 0。那么我们知道 $x_2 = 0 + \epsilon_2$ 。因此, x_2 的方差 $= \text{Var}(\epsilon_2) = \sigma^2$ 。现在, $x_3 = x_2 + \epsilon_3 = \epsilon_2 + \epsilon_3$ 。因为每一期的误差项与其他期的误差项之间被假定为不相关的, 所以 x_3 的方差 $= \text{Var}(\epsilon_2) + \text{Var}(\epsilon_3) = 2\sigma^2$ 。以相同的方法, 我们可以得到对于任一期 t , x_t 的方差 $= (t-1)\sigma^2$ 。但是, 这意味着随着 t 的增加, x_t 的方差将不断增长而没有上界: 它会达到正无穷。缺乏上界反过来意味着随机游走不是协方差平稳的时间序列, 因为一个协方差平稳的时间序列必须具有有限的方差。

这些问题所隐含的实际意义是什么呢? 我们不能使用标准的回归分析来处理一个服从随机游走的时间序列。然而, 如果我们怀疑一个时间序列是随机游走的, 那么我们可以试图将该数据转化为协方差平稳的时间序列。用统计学的术语来说, 我们可以将其差分。

我们将一个时间序列进行差分, 就是产生一个新的时间序列, 比如说 y_t , 其每一期的值等于 x_t 和 x_{t-1} 之差。这一转换被称为一阶差分 (first-differencing), 因为它将前一期时间序列的值从当期的时间序列值中减去。有时, x_t 的一阶差分被写成 $\Delta x_t = x_t - x_{t-1}$ 。注意对式 (10-8) 中的随机游走进行一阶差分后得到

$$y_t = x_t - x_{t-1} = \epsilon_t, E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2, E(\epsilon_t \epsilon_s) = 0 \quad \text{如果 } t \neq s$$

⊖ 式 (10-8) 加上了一个非 0 的截距 (如同之后给出的式 (10-9)), 该截距有时也被称为一个随机游走的漂移项。

ϵ_t 的期望值为 0。因此，在时刻 $t-1$ 对于 y_t 的最佳预测为 0。这表示最佳预测认为当前时间序列的值 x_{t-1} 将不会发生变化。

一阶差分后的变量 y_t 是协方差平稳的。为什么呢？首先，注意到这个模型 ($y_t = \epsilon_t$) 是一个 AR (1) 模型，其中 $b_0 = 0$, $b_1 = 0$ 。我们可以计算一阶差分后模型的均值回复水平，其值为 $b_0/(1 - b_1) = 0/1 = 0$ 。因此，一阶差分后的随机游走具有的均值回复水平为 0。同时注意到每一期 y_t 的方差为 $\text{Var}(\epsilon_t) = \sigma^2$ 。因为每一期 y_t 的均值和方差都为有限常数，所以 y_t 是协方差平稳的时间序列，而我们可以使用线性回归来对其建模。[⊖]当然，对一阶差分后的序列用 AR (1) 模型建模并不能帮助我们预测未来，因为 $b_0 = 0$, $b_1 = 0$ 。事实上，我们简单地得到原始的时间序列是一个随机游走的结论。

假使我们试图对于一个随机游走的时间序列用 AR (1) 模型进行估计，那么我们的统计结果将是不正确的，因为 AR 模型无法用来估计随机游走或者其他任何不平稳的时间序列。下面的例子就用汇率说明了这个问题。

例 10-10 日元/美元汇率

金融分析师经常假设汇率服从随机游走模型。考虑一个关于日元/美元汇率的 AR (1) 模型。表 10-9 给出了使用 1975 年 1 月到 2002 年 12 月的月末观测值所得到的模型估计结果。

表 10-9 日元/美元汇率：AR (1) 模型月底观测值 (1975 年 1 月~2002 年 12 月)

回归统计量			
R 平方		0.991 4	
标准误		5.900 6	
观测数		336	
杜宾—沃特森		1.849 2	
	系数	标准误	t 统计量
截距	1.022 3	0.926 8	1.109 2
滞后 1 期	0.991 0	0.005 0	196.451 7
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	0.070 6	0.054 6	1.293 0
2	0.036 4	0.054 6	0.666 7
3	0.086 4	0.054 6	1.582 4
4	0.056 6	0.054 6	1.036 6

资料来源：U. S. Federal Reserve Board of Governors.

表 10-9 的结果表明日元/美元汇率是一个随机游走，因为所估计出的截距没有表现出显著不同于 0，而估计出的汇率一阶滞后项的系数也非常接近于 1。我们可以使用表 10-9 中的 t 统计量来检验汇率是否服从随机游走模型吗？不幸的是，我们不能这么做，因为如果模型被用来对于随机游走进行估计，那么 AR 模型的标准误将是无效的（记住，随机游走不是协方差平稳的）。如果汇率事实上确实是一个随机游走，那么我们根据错误的统计检验就可能得到一个错误的结论，然后错误地进行投资。我们可以使用在下一节中给出的方法来对于时间序列是否服从随机游走进行检验。

⊖ 所有的协方差都是有限的，这主要是由于两个原因：首先，方差是有限的；其次，一个时间序列和其历史值的协方差不会大于该时间序列的方差。

假设汇率正如我们所怀疑的，是一个随机游走。如果真是这样的话，那么一阶差分序列， $y_t = x_t - x_{t-1}$ ，将是协方差平稳的。我们在表 10-10 中给出了 $y_t = b_0 + b_1 y_{t-1} + \epsilon_t$ 的估计结果。如果汇率是一个随机游走，那么 $b_0 = 0$ ， $b_1 = 0$ ，并且误差项将是序列不相关的。

表 10-10 一阶差分后的日元/美元汇率：AR（1）模型月底观测值（1975 年 1 月~2002 年 12 月）

回归统计量			
R 平方		0.005 3	
标准误		5.913 3	
观测数		336	
杜宾—沃特森		1.998 0	
	系数	标准误	t 统计量
截距	-0.496 3	-0.496 3	-1.530 1
滞后 1 期	0.072 6	0.054 7	1.328 2
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	-0.004 5	0.054 6	-0.082 4
2	0.025 9	0.054 6	0.474 4
3	0.080 7	0.054 6	1.478 0
4	0.048 8	0.054 6	0.893 8

资料来源：U. S. Federal Reserve Board of Governors.

在表 10-10 中，截距和汇率一阶差分后的一阶滞后项的系数都不显著异于 0，而且残差自相关系数都不异于 0。[⊖] 这些结果与日元/美元汇率服从一个随机游走的假设是一致的。

我们已经得出结论：差分后的回归模型是我们所要选择的模型。现在我们能够看到如果我们基于模型的 R^2 进行比较的话，那么我们将会严重地被误导。在表 10-9 中， R^2 为 0.9914，而在表 10-10 中， R^2 为 0.005 3。如果我们只是认为表 10-10 的模型是我们应该使用的，那么可能的结果会怎么样的呢？在表 10-9 中， R^2 度量了上一期汇率预测下一期汇率的好坏程度。如果汇率是一个随机游走，那么其当前值将是非常好的对于下一期汇率的预测值，因此， R^2 将是非常高的。同时，如果汇率是一个随机游走，那么汇率的变化应该完全无法被预测。表 10-10 对于上个月到下个月的汇率变化是否可以由前个月的汇率变化所预测给出了估计。如果我们无法预测，那么表 10-10 的 R^2 应该很低。事实上，它确实很低（0.005 3）。这一比较提供了我们一个很好的例子，该例子说明我们无法只通过比较两个模型的 R^2 来选择哪个模型是正确的。

汇率是一个随机游走，而根据定义，一个随机游走的变化是无法预测的。因此，我们无法通过预测汇率变化的投资策略来获利。

到目前为止，我们仅讨论了随机游走；即不带有漂移项的随机游走。在一个带有漂移项的随机游走中，下一期时间序列的最佳预测值就是当前值。然而，带有漂移项的随机游走应该在每一期以一个常数增加或者减少。描述带有漂移项的随机游走的等式是 AR（1）模型的一个特例：

$$x_t = b_0 + b_1 x_{t-1} + \epsilon_t$$
$$b_1 = 1, b_0 \neq 0, \text{或者}$$

⊖ 参见 Greene（2003）中对于回归系数都等于 0 的联合假设的检验。

$$x_t = b_0 + x_{t-1} + \epsilon_t, E(\epsilon_t) = 0 \quad (10-9)$$

带有漂移项的随机游走，相对于 $b_0 = 0$ 的简单随机游走，其 $b_0 \neq 0$ 。

我们已经看到 $b_1 = 1$ 暗示了一个无法定义的均值回复水平，因此其是不平稳的。所以，我们无法使用 AR 模型来分析带有漂移项的随机游走的时间序列，除非我们将时间序列进行一阶差分，将其转化。如果我们将式 (10-9) 一阶差分，那么其结果为 $y_t = x_t - x_{t-1}$, $y_t = b_0 + \epsilon_t$, $b_0 \neq 0$ 。

10.5.2 非平稳数据的单位根检验

在这一节中，我们将讨论如何使用随机游走的概念来判断一个时间序列是否是协方差平稳的。这一方法聚焦于 AR (1) 模型中带有漂移项的随机游走的斜率系数，而不是我们之前讨论的传统的自相关系数的方法。

时间序列不同滞后期自相关系数的检验是我们熟知的用来推断一个时间序列是否是平稳的处理方法。通常，对于一个平稳的时间序列，要么其所有滞后期的自相关系数都在统计上不显著为 0，要么其自相关系数会随着滞后期数的增加，迅速地下降到 0。相反，一个不平稳的时间序列则不会表现出这些特点。然而，这种方法与当前更加流行的检验非平稳性的方法（单位根的迪克—福勒 (Dickey-Fuller) 检验）相比，判断标准不够明确。

我们可以在 AR (1) 模型中解释什么是所谓的单位根问题。如果一个时间序列来自于一个 AR (1) 模型，那么要使得其是协方差平稳的，滞后项系数 b_1 的绝对值，必须小于 1.0。如果滞后项系数的绝对值大于等于 1.0，那么我们就不能使用 AR (1) 模型的统计结果，因为该时间序列不是协方差平稳的。如果滞后项系数等于 1.0，那么时间序列就具有一个单位根 (unit root)：它是一个随机游走而且不是协方差平稳的。^①根据定义，所有随机游走，包括漂移项和不包括漂移项的随机游走，都具有单位根。

我们如何检验一个时间序列是否存在单位根呢？如果我们认为一个时间序列 x_t 是一个带有漂移项的随机游走，那么我们可能会利用线性回归来估计 AR (1) 模型 $x_t = b_0 + b_1 x_{t-1} + \epsilon_t$ ，并对于 $b_1 = 1$ 的假设进行 t 检验。不幸的是，如果 $b_1 = 1$ ，那么 x_t 不是协方差平稳的，而估计出的系数 $\hat{b}_1 = 1$ ，实际上并不服从 t 分布；因此， t 检验是无效的。

迪克 (Dickey) 和福勒 (Fuller) (1979) 提出了关于 AR (1) 模型的转化版本 $x_t = b_0 + b_1 x_{t-1} + \epsilon_t$ 的一个基于回归的单位根检验方法。将 AR (1) 模型的两边同时减去 x_{t-1} ，得到

$$x_t - x_{t-1} = b_0 + (b_1 - 1)x_{t-1} + \epsilon_t$$

或者

$$x_t - x_{t-1} = b_0 + g_1 x_{t-1} + \epsilon_t, E(\epsilon_t) = 0 \quad (10-10)$$

式中， $g_1 = (b_1 - 1)$ 。如果 $b_1 = 1$ ，那么 $g_1 = 0$ ，因此，检验 $g_1 = 0$ 也就是检验 $b_1 = 1$ 。如果在 AR (1) 模型中存在一个单位根，那么回归中的 g_1 将为 0，其中因变量是时间序列的一阶差分，而自变量是时间序列的一阶滞后项。迪克—福勒检验的原假设为 $H_0: g_1 = 0$ ——即该时间序列具有一个单位根，而且是不平稳的。为了进行该检验，我们采用传统的方法来计算 g_1 的 t 统计量，但是对于其临界值，我们使用由迪克和福勒所计算的修正临界值，而不是使用传统 t 检验的临界值；修正的临界值的绝对值大于传统的临界值。许多软件包都包含迪克—福勒检验。^②

① 当 b_1 的绝对值大于 1 时，我们就认为存在一个使时间序列迅速增长的根。更多内容参见 Diebold (2004)。

② Dickey 和 Fuller 建立了三种不同模型（随机游走，带有漂移项的随机游走，或者带有漂移项和趋势项的随机游走）假定下对于 $g_1 = 0$ 假设的检验。对于这三种模型的迪克—福勒检验的临界值是不同的。关于该问题的更多内容，参见 Greene (2003) 或 Hamilton (1994)。

例 10-11 英特尔公司季度销售额 (1)

先前, 我们得到了这样的结论: 我们无法只利用时间趋势线 (如例 10-3 所示) 来对于英特尔公司季度销售额的对数进行建模。回忆一下, 对数—线性回归中的杜宾—沃特森统计量使得我们拒绝了回归中的误差是序列不相关的原假设。而假如分析师决定利用 AR (1) 模型来对于英特尔公司季度销售额的对数进行建模。那么, 他会使用 $\ln \text{销售额}_t = b_0 + b_1 \ln \text{销售额}_{t-1} + \epsilon_t$ 。

在分析师估计该回归之前, 他应该使用迪克—福勒检验来确定英特尔公司季度销售额的对数序列是否存在一个单位根。如果他使用 1985 年第一季度到 1999 年第四季度的英特尔公司销售额的季度数据样本, 对于每个观测值取自然对数, 并计算迪克—福勒检验的 t 统计量, 该统计量的值可能使得他无法拒绝英特尔公司季度销售额的对数序列存在一个单位根的原假设。

如果一个时间序列可能具有一个单位根, 那么我们应该如何来对其建模呢? 一种通常可行的方法是将该时间序列进行一阶差分 (正如之前所讨论的), 并尝试对一阶差分序列用自回归时间序列模型来建模。下面的例子给出了这一方法。

例 10-12 英特尔公司季度销售额 (2)

假如你相信通过对于时间序列图像的观测, 英特尔公司季度销售额的对数不是协方差平稳的 (它具有一个单位根)。因此, 你构造了一个新的序列 y_t , 其是英特尔公司季度销售额的对数序列的一阶差分。图 10-10 反映了这一序列。

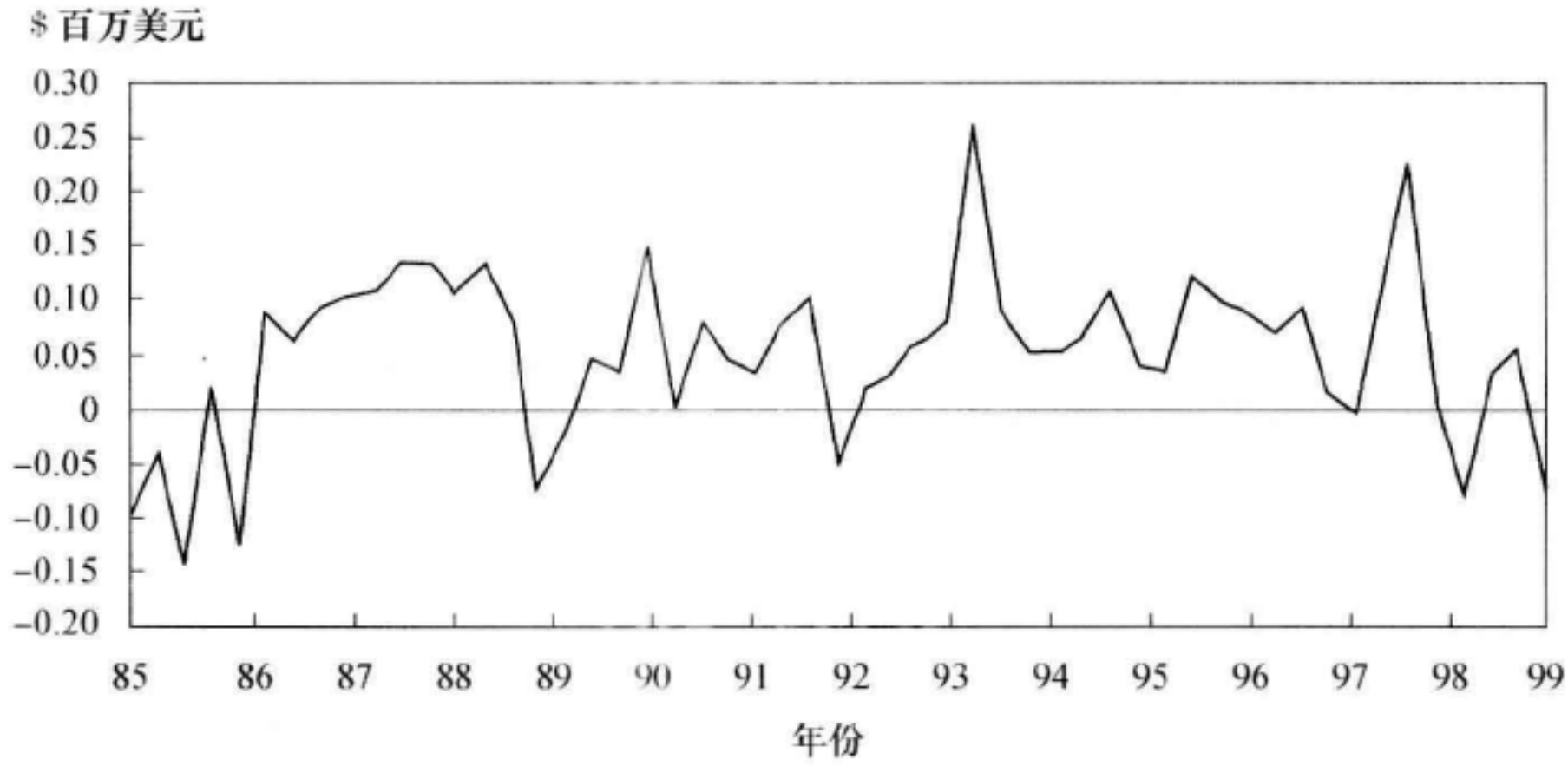


图 10-10 对数差分, 英特尔公司季度销售额

资料来源: Compustat.

如果你将图 10-10 和图 10-6 以及图 10-8 进行比较, 你将看到英特尔公司季度销售额的对数序列的一阶差分消除了英特尔公司销售额序列和英特尔公司销售额对数序列中所反映出的很强的向上趋势。因为一阶差分序列不具有很强的趋势, 我们能更好地假设差分序列是协方差平稳的, 而不是假设英特尔公司销售额序列或者英特尔公司销售额对数序列是协方差平稳的时间序列。

现在假设你决定利用 AR (1) 模型对新序列进行建模。你会使用 $\ln(\text{销售额}_t) - \ln(\text{销售额}_{t-1}) = b_0 + b_1[\ln(\text{销售额}_{t-1}) - \ln(\text{销售额}_{t-2})] + \epsilon_t$ 。表 10-11 给出了这个回归的估计结果。

表 10-11 对数差分销售额序列: AR (1) 模型, 英特尔公司季度观测值 (1985 年 1 月 ~ 1999 年 12 月)

回归统计量			
R 平方	0.094 6		
标准误	0.075 8		
观测数	60		
杜宾—沃特森	1.970 9		
	系数	标准误	t 统计量
截距	0.035 2	0.011 4	3.087 5
滞后 1 期	0.306 4	0.124 4	2.462 0
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	-0.014 0	0.129 1	-0.108 8
2	-0.085 5	0.129 1	-0.662 4
3	-0.058 2	0.129 1	-0.450 6
4	0.212 5	0.129 1	1.646 3

资料来源: Compnstat.

表 10-11 的下半部分表明该模型的前四阶残差的自相关系数很小。在 60 个观测值和 2 个参数的情况下, 该模型具有 58 个自由度。在 0.05 显著性水平下, 该模型 t 统计量的临界值约为 2.0。这些自相关系数的 t 统计量的绝对值没有一个大于 2.0。因此, 我们无法拒绝每个自相关系数等于 0 的原假设。我们认为残差没有表现出明显的自相关系数。

该结果表明模型是设定良好的, 并且我们可以使用其所估计的结果。新的一阶差分序列的一阶滞后项的截距 ($\hat{b}_0 = 0.035 2$) 和系数 ($\hat{b}_1 = 0.306 4$) 都是统计上显著的。我们如何来解释模型中的这些估计系数呢? 截距的值 (0.0352) 表明如果销售额在当前季度没有发生改变 ($y_t = \ln \text{销售额}_t - \ln \text{销售额}_{t-1} = 0$), 销售额将在下个季度增长 3.52%。[⊖] 然而, 如果销售额在当前季度发生了改变, 该模型预测在该季度的销售额将增长 3.52% 加上 0.306 4 乘以销售额的增长。

假如我们想要使用该模型, 在 1999 年第四季度末预测 2000 年第一季度的英特尔公司的销售额。我们令 t 表示 1999 年第四季度, 因此, $t-1$ 表示 1999 年第三季度, 而 $t+1$ 表示 2000 年第一季度。于是, 我们将计算 $\hat{y}_{t+1} = 0.035 2 + 306 4 y_t$ 。为了计算 $\hat{y}_{t+1}$, 我们需要知道 $y_t = \ln \text{销售额}_t - \ln \text{销售额}_{t-1}$ 。在 1999 年第三季度, 英特尔公司的销售额为 7 328 百万美元, 因此, $\ln(\text{销售额}_{t-1}) = \ln 7\,328 = 8.899 5$ 。在 1999 年第四季度, 英特尔公司销售额为 8 212 百万美元, 所以 $\ln(\text{销售额}_t) = \ln 8\,212 = 9.013 4$ 。故 $y_t = 9.013 4 - 8.899 5 = 0.113 9$ 。因此, $\hat{y}_{t+1} = 0.035 2 + 306 4 (0.113 9) = 0.070 1$ 。如果 $\hat{y}_{t+1} = 0.070 1$, 那么 $0.070 1 = \ln(\text{销售额}_{t+1}) - \ln(\text{销售额}_t) = \ln(\text{销售额}_{t+1}/\text{销售额}_t)$ 。如果我们对于该等式两边同取指数, 该结果为

⊖ 注意 3.52% 是指数增长率, 而不是 $[(\text{当前季度销售额}/\text{前一季度销售额}) - 1]$ 。这两种计算增长率的方法之间的差别通常很小。

$$e^{0.0701} = (\text{销售额}_{t+1} / \text{销售额}_t)$$

$$\begin{aligned}\text{销售额}_{t+1} &= \text{销售额}_t e^{0.0701} = \$ 8\,212 \text{ 百万美元} \times 1.0726 \\ &= \$ 8\,808 \text{ 百万美元}\end{aligned}$$

因此,在1999年第四季度,该模型将预测2000年第一季度的英特尔公司销售额将为8808百万美元。该销售额预测可能会影响到我们在那个时点购买英特尔公司股票的决策。

10.6 移动平均时间序列模型

至今,我们所使用的许多预测模型都是自回归模型。因为大部分金融时间序列具有自回归过程的性质,自回归时间序列模型很可能是在金融预测中最常使用的时间序列模型。然而,一些金融时间序列似乎更接近服从于另一种被称为移动平均的时间序列模型。例如,正如我们在之后将会看到的,比起自回归过程,标准普尔500指数的收益率可能更适合用移动平均过程来建模。

在本节中,我们将介绍一些移动平均模型的基本概念,并使得你可以在给出这一模型时,能够提出相关的问题。我们首先讨论如何利用移动平均来平滑历史值,然后讨论如何利用移动平均模型来预测一个时间序列。即使这两种方法都包含了“移动平均”这一词汇,但是它们的含义是不同的。

10.6.1 用 n 期移动平均平滑历史数据

假设你正在分析一家公司历史销售额的长期趋势。为了关注于趋势,你可能会发现通过将销售额的时间序列进行平滑,以消除短期波动或者噪声是非常有用的。一种消除一个时间序列中时间段与时间段之间波动的方法是 n 期移动平均 (n -period moving average)。一个关于当期和过去 $n-1$ 期时间序列 x_t 值的 n 期移动平均计算为

$$\frac{x_t + x_{t-1} + \cdots x_{t-(n-1)}}{n} \quad (10-11)$$

下面的例子展示了如何计算英特尔公司季度销售额的一个移动平均。

例 10-13 英特尔公司季度销售额 (3)

假如我们想要计算1999年第四季度末英特尔公司销售额的四季度移动平均。前四个季度英特尔公司销售额分别为1999年第一季度7103百万美元;1999年第二季度6746百万美元;1999年第三季度7328百万美元以及1999年第四季度8212百万美元。因此,2000年第一季度所计算的四季度销售额的移动平均为 $(7\,103 + 6\,746 + 7\,328 + 8\,212)/4 = 7\,347.25$ 百万美元。

我们通常绘制一个具有剧烈波动的时间序列的移动平均序列来帮助我们看清数据的某种模式。图10-11给出了1972年1月到2000年12月的月度实际的(经过通货膨胀调整的)美国零售销售额,同时,也给出了数据12个月的移动平均值。[⊖]

⊖ 一个12个月的移动平均是一个时间序列最近12个月的平均值。虽然样本区间从1972年开始,但是1971年的数据被用来计算1972年每个月的12个月的移动平均。

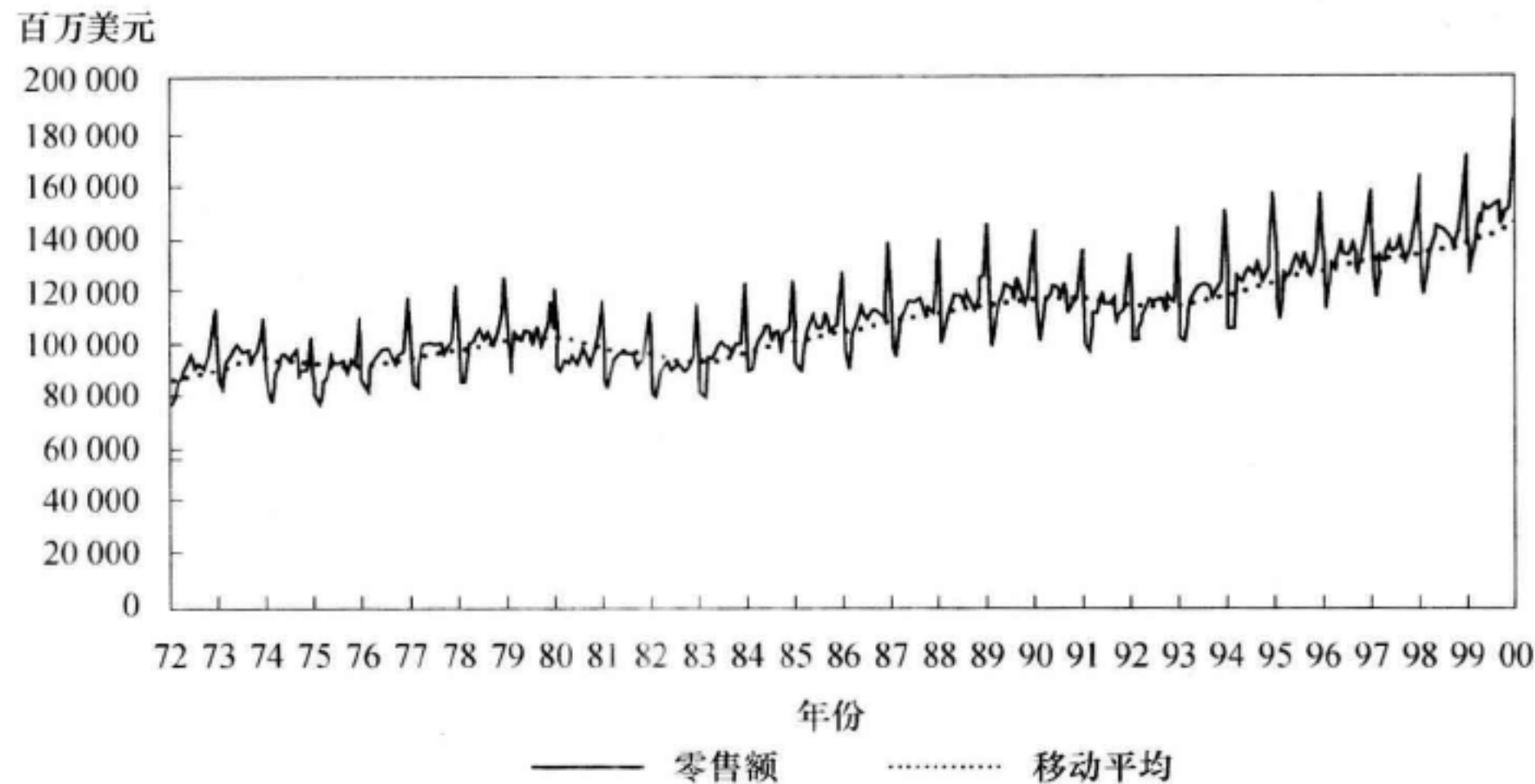


图 10-11 月度美国实际零售额和 12 个月零售额的移动平均

资料来源：U. S. Department of Commerce, Census Bureau.

正如图 10-11 所示，每年零售额具有一个很高的高峰（12 月），同时紧接着一个很低的低谷（1 月）。因为数据中极端的季节性，一个 12 个月的移动平均能帮助我们聚焦于零售额的长期变动情况而不是季节性的波动情况。注意移动平均不具有原始零售额数据中明显的季节性波动。相反，例如，从 1985 ~ 1990 年零售额的移动平均会稳步增长，而 1990 ~ 1993 年会下降，然后在稳步上升。12 个月的移动平均可以比观测时间序列本身更容易看到该时间序列的变动趋势。

图 10-12 给出了美国月度原油价格以及 12 个月原油价格的移动平均值。虽然这些数据不具有如同图 10-11 中零售额数据所给出的明显规则的季节性，但是移动平均消除了月度原油价格的波动，从而显示出了其长期变动情况。

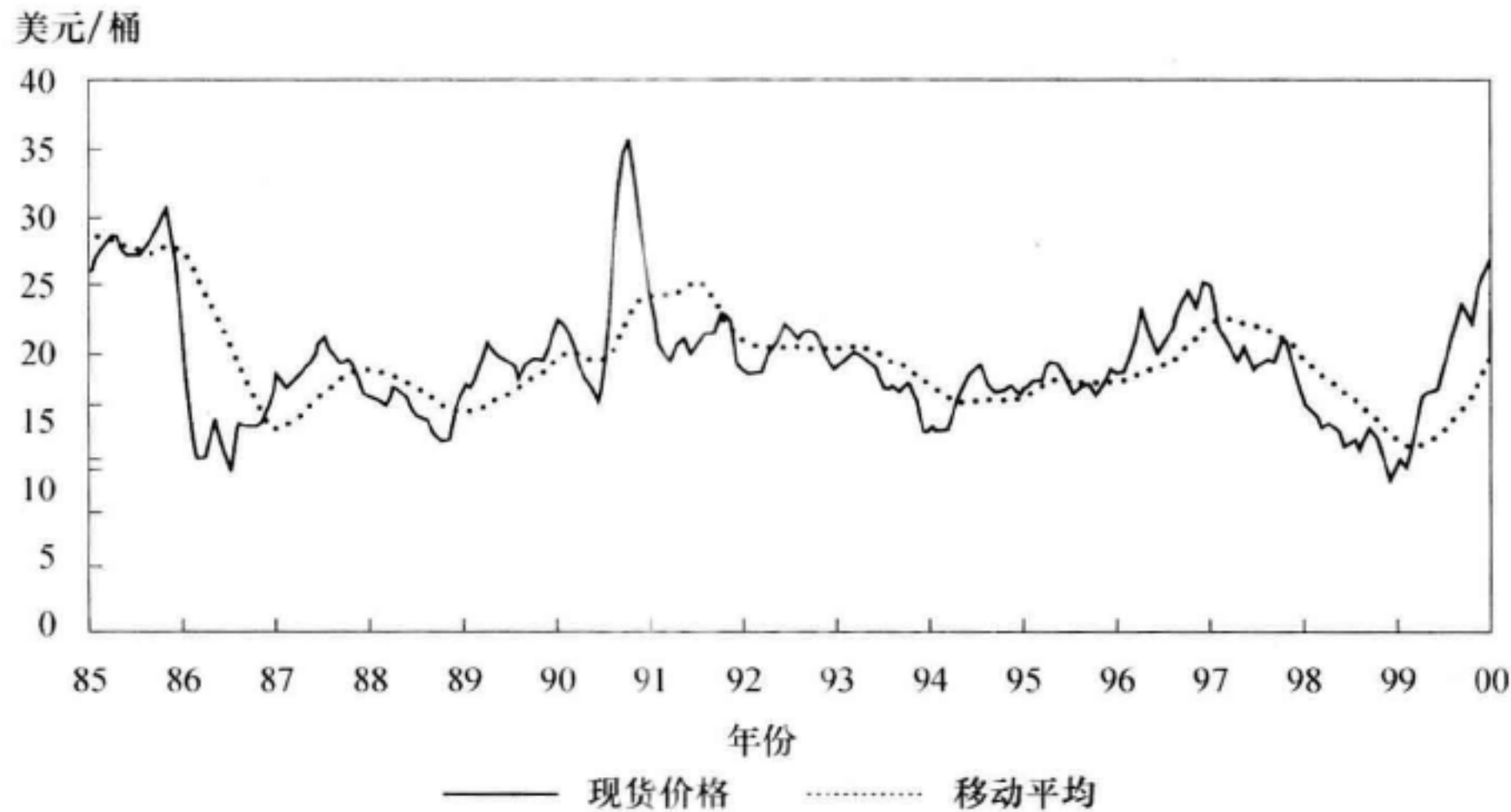


图 10-12 月度原油价格和 12 个月价格的移动平均

资料来源：Dow Jones Energy Service.

图 10-12 同样也反映出移动平均的一个缺点：它总是滞后与实际数据。例如，当原油价格在 1985 年后期大幅下跌并保持在相对较低位置时，移动平均只是在缓步下降。而当原油价格在

1999 年迅速上升时, 移动平均同样也滞后了。因此, 一个简单的关于当前历史值的移动平均, 虽然经常在平滑一个时间序列中非常有用, 但是可能并不是一个对于未来最佳的预测指标。该问题的主要原因在于一个简单的移动平均对于移动平均中的所有区间给予相同的权重。为了预测一个时间序列的将来值, 通常最好需要使用更加复杂的移动平均时间序列模型。我们将在下面讨论这种模型。

10.6.2 用移动平均时间序列模型来进行预测

假如一个时间序列 x_t 是与下列模型相一致的:

$$x_t = \epsilon_t + \theta\epsilon_{t-1}, E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2, E(\epsilon_t\epsilon_s) = 0 \quad \text{对于 } t \neq s \quad (10-12)$$

该等式被称为一阶移动平均模型, 或者简称为 MA (1) 模型。西塔 (θ) 是 MA (1) 模型的系数。[⊖]

式 (10-12) 是一个移动平均模型, 因为在每一期, x_t 是一个关于 ϵ_t 和 ϵ_{t-1} (两个不相关随机变量, 且每个变量的期望值为 0。) 的移动平均。不同于式 (10-11) 的简单移动平均模型, 该移动平均模型对于移动平均中的两项赋予了不同的权重 (ϵ_t 给予 1, 而 ϵ_{t-1} 给予 θ)。

我们可以通过观察一个时间序列的自相关系数来确定 x_t 是否只与其之前和之后的值相关, 从而判断该时间序列是否满足一个 MA (1) 模型。首先, 我们检查式 (10-12) 中 x_t 的方差以及其前两阶自相关系数。因为 x_t 在所有期的期望值为 0, 并且 ϵ_t 与其自身的历史值是不相关的, 所以第一个自相关系数不等于 0, 但是第二个以及更高阶的自相关系数都等于 0。进一步的分析表明在 MA (1) 模型中, 除了第一阶的自相关系数, 其余所有的自相关系数都等于 0。因此, 对于一个 MA (1) 过程, 任何 x_t 的值都与 x_{t-1} 和 x_{t+1} 有关, 但是与其他时间序列值无关; 我们可以认为一个 MA (1) 模型具有一期的记忆。

当然, 一个 MA (1) 模型不是最复杂的移动平均模型。一个 q 阶的移动平均模型, 记为 MA (q), 其具滞后项具有不同的权重, 它可以写为

$$x_t = \epsilon_t + \theta_1\epsilon_{t-1} + \cdots + \theta_q\epsilon_{t-q}, E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2 \quad (10-13)$$

$$E(\epsilon_t\epsilon_s) = 0 \quad \text{对于 } t \neq s$$

那么我们如何说明 MA (q) 模型是否适合一个时间序列呢? 我们可以检查其自相关系数。对于 MA (q) 模型, 前 q 个自相关系数将显著不为 0, 而之后所有的自相关系数将等于 0; 一个 MA (q) 模型具有 q 期的记忆。这个结果对于选择 MA 模型中正确的 q 值非常重要。我们在上面对于其特例 $q=1$, 即除了第一阶之外所有自相关系数等于 0 的 MA (1) 模型中已经讨论了该结果。

那么我们如何来区分自回归时间序列和移动平均时间序列呢? 我们要再一次通过检查时间序列本身的自相关系数来回答该问题。大部分自回归时间序列的自相关系数开始很大, 然后逐渐下降, 而 MA (q) 时间序列的自相关系数在前 q 阶自相关系数之后会突然降到 0。我们不可能事先知道一个时间序列是自回归序列还是移动平均序列。因此, 自相关系数给予我们关于如何对时间序列建模的最佳提示。然而, 大部分时间序列还是以自回归模型建模的效果为最好。

⊖ 注意正如 6.1 节中所述, 一个移动平均时间按序列模型与简单的移动平均非常不同。简单的移动平均是基于一个时间序列的观测值。而在一个移动平均时间序列模型中, 我们从不直接观测 ϵ_t 或者其他任何 ϵ_{t-j} , 但是如我们下面所讨论的, 我们能推断出一个特定的移动平均模型是如何反映一个时间序列的某种序列相关模式的。

例 10-14 月度标准普尔 500 收益率的时间序列模型

标准普尔月度收益率是自相关的吗？如果是的话，我们就可能设计出一种投资策略来利用该自相关性。对于标准普尔 500 月度收益率最合适的时间序列模型是什么？

表 10-12 利用 1991 年 1 月到 2002 年 12 月的月度数据给出了标准普尔 500 收益率的前六阶自相关系数。注意到所有的自相关系数都很小。它们是否是显著的？在 144 个观测值 0.05 显著性水平下，该模型 t 值的临界值约为 1.98。没有一个自相关系数的 t 统计量的绝对值超过该临界值 1.98。因此，我们无法拒绝这些自相关系数显著不同于 0 的原假设。

表 10-12 年化的月度标准普尔 500 收益率 (1991 年 1 月~2002 年 12 月)

自相关系数			
滞后期	自相关系数	标准误	t 统计量
1	-0.009 0	0.083 3	-0.108 3
2	-0.020 7	0.083 3	-0.248 1
3	0.002 0	0.083 3	0.024 0
4	-0.073 0	0.083 3	-0.875 6
5	0.114 3	0.083 3	1.371 7
6	-0.000 7	0.083 3	-0.008 2
观测值	144		

资料来源：Ibbotson Associates.

如果标准普尔 500 的收益率是一个 $MA(q)$ 时间序列，那么前 q 阶自相关系数将显著不同于 0。然而，没有一个自相关系数是显著的，所以标准普尔 500 的收益率似乎来自 $MA(0)$ 时间序列。一个 $MA(0)$ 时间序列是指我们让该模型的均值非 0，并采用下列形式：[⊖]

$$x_t = \mu + \epsilon_t, E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2, E(\epsilon_t \epsilon_s) = 0 \quad \text{对于 } t \neq s \tag{10-14}$$

其意味着该时间序列是不可预测的。该结果不应该令我们太惊讶，正如大部分研究所表明的，股票指数的短期收益率是非常难以预测的。

我们能从该例子中看到对于自相关系数的检验是如何帮助我们在 AR 和 MA 模型之间进行选择的。如果标准普尔 500 的收益率是来自于 AR(1) 时间序列，那么其一阶自相关系数将显著不同于 0，而其他自相关系数将逐渐下降。然而，这里甚至一阶自相关系数也不是显著异于 0 的。因此，我们可以确信标准普尔 500 的收益率不来自于 AR(1) 模型，或者任何更高阶的 AR 模型（由于相同的原因）。这个发现是与我们标准普尔 500 序列是 $MA(0)$ 的结论相一致的。

10.7 时间序列模型中的季节性

在我们对本章中的时间序列模型结果进行分析时，我们会面临一些问题。一个常见的问题是

⊖ 基于投资理论和实际证据，我们预计标准普尔 500 的平均月度收益率为正 ($\mu > 0$)。我们也能通过加上一个常数项 μ 来一般化刻画 $MA(q)$ 时间序列的式 (10-13)。将常数项包含在一个移动平均模型中并没有改变对于时间序列方差和自协方差的表述。许多对于弱式有效市场的早期研究都是采用式 (10-14) 来对于股票收益率建模的。参见 Garbade (1982)。

显著的季节性，即序列在每年中反映出规则的变动情况。乍一看，季节性可能会使我们无法使用自回归时间序列模型。毕竟，自相关系数将根据季节不同而不同。然而，该问题通常可以通过在自回归模型中使用季节滞后项而得到解决。

季节性滞后项一般是包含在自回归模型中的额外项，其等于当期之前一年的时间序列值。例如，假设我们对于某个季度时间序列用 AR (1) 模型建模， $x_t = b_0 + b_1 x_{t-1} + \epsilon_t$ 。如果一个时间序列具有明显的季节性，那么该模型将不是正确设定的。季节性可以被很容易地检测出，因为季节性的误差项自相关系数（对季度数据来说，是第四阶自相关系数）将显著不等于 0。假如该季度模型具有显著的季节性。在这个例子中，我们可能在自回归模型中包含一个季节性滞后项，并估计

$$x_t = b_0 + b_1 x_{t-1} + b_2 x_{t-4} + \epsilon_t \tag{10-15}$$

来检验是否包含季节性滞后项的模型会消除误差项中统计上显著的自相关系数。

在例 10-15 和例 10-16 中，我们将给出如何在一个时间序列模型中检验和调整季节性效应。我们也将给出如何使用一个带有季节性滞后项的自回归模型来计算预测值。

例 10-15 美敦力 (Medtronic) 公司销售额序列中季节性效应

我们想要预测美敦力公司销售额。基于本章中先前的结果，我们确定销售额对数的一阶差分很可能是协方差平稳的。利用 1985 年第一季度到 2001 年最后一季度的季度数据，我们使用普通最小二乘法对于一阶差分数据的 AR (1) 模型进行估计。我们估计下面的等式： $(\ln \text{销售额}_t - \ln \text{销售额}_{t-1}) = b_0 + b_1 (\ln \text{销售额}_{t-1} - \ln \text{销售额}_{t-2}) + \epsilon_t$ 。表 10-13 给出了回归结果。

表 10-13 对数差分销售额：AR (1) 模型美敦力公司季度观测值 (1985 年 1 月 ~ 2001 年 12 月)

回归统计量			
R 平方	0.1619		
标准误	0.0693		
观测数	68		
杜宾—沃特森	2.0588		
	系数	标准误	t 统计量
截距	0.0597	0.0091	6.5411
滞后 1 期	-0.4026	0.1128	-3.5704
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	-0.0299	0.1213	-0.2463
2	-0.1950	0.1213	-1.6077
3	-0.1138	0.1213	-0.9381
4	0.4072	0.1213	3.3581

资料来源：Compustat.

表 10-13 中首先要注意的是残差具有很强的季节性自相关系数。该表最下面的部分表明第四阶自相关系数具有值 0.4072，而其 t 统计量为 3.36。在 68 个观测值和 2 个参数的情况下，

该模型具有 66 个自由度。[⊖]在 0.05 显著性水平下, t 统计量的临界值约为 2.0。给定 t 统计量的这个值, 我们必须拒绝第四阶自回归系数等于 0 的原假设, 因为 t 统计量对于临界值 2.0。在这个模型中, 第四阶自相关系数是季度性的自相关系数, 因为这个 AR (1) 模型是用季度数据估计的。表 10-13 给出了很强的统计上显著的季度自相关系数, 其在对于带有很强季节效应的时间序列用不考虑季节性效应的模型建模时发生。因此, AR (1) 模型是设定错误的。我们不应该使用该模型来进行预测。

假如我们由于季节性的自相关系数, 决定使用带有一个季节性滞后项的自回归模型。我们对于季度数据建模, 因此我们估计式 (10-15): $(\ln \text{销售额}_t - \ln \text{销售额}_{t-1}) = b_0 + b_1 (\ln \text{销售额}_{t-1} - \ln \text{销售额}_{t-2}) + b_2 (\ln \text{销售额}_{t-4} - \ln \text{销售额}_{t-5}) + \epsilon_t$ 。该等式的估计结果反映在表 10-14 中。

表 10-14 对数差分销售额: 带有季节滞后项的 AR (1) 模型, 美敦力公司
季度观测值 (1985 年 1 月~2001 年 12 月)

回归统计量			
R 平方	0.3219		
标准误	0.0580		
观测数	68		
杜宾—沃特森	2.0208		
	系数	标准误	t 统计量
截距	0.0403	0.0096	4.1855
滞后 1 期	-0.2955	0.1058	-2.7927
滞后 4 期	0.3896	0.0995	3.9159
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	-0.0108	0.1213	-0.0889
2	-0.0957	0.1213	-0.7889
3	0.0075	0.1213	0.0621
4	-0.0340	0.1213	-0.2801

资料来源: Compustat.

注意到表 10-14 的底部给出了残差的自相关系数。当我们在回归中包含一个季节滞后项之后, 前四阶自相关系数的 t 统计量没有一个保持显著。

现在我们了解到该模型的残差不具有显著的序列相关性, 所以我们假设该模型是正确设定的。我们如何来解释该模型中的系数呢? 为了预测美敦力公司当前季度的销售增长率, 我们需要知道两个数值: 前一季度的销售增长率和四个季度之前的销售增长率。如果销售额在这两个季度保持不变, 那么表 10-14 中的模型预测销售额将在当前季度增长 0.0403 (4.03%)。如果上个季度销售额增长 1%, 而四个季度之前的销售额增长 2%, 那么该模型预测当前季度的销售增长率将为 $0.0403 - 0.2955 (0.01) + 0.3896 (0.02) = 0.0451$ 或者 4.51%。[⊗]同样注意到带有季节滞后项模型的 R^2 (表 10-14 中的 0.3219) 几乎是不带有季节滞后项模型 (表 10-13 中的 0.1619) 的两倍。因此, 带有季节滞后项的模型能够更好地解释该数据。

⊖ 虽然样本区间从 1985 年开始, 但是我们使用之前的观测值作为滞后项。不然, 该模型将具有更少的自由度, 因为样本大小将会随着滞后项数量的增加而降低。
⊗ 注意所有这些增长率是指数增长率。

例 10-16 零售额增长率

我们想要预测美国月度零售额的增长率，来决定是否推荐折扣商店（discount store）的股票。我们决定使用没有经过季节性调整，但是经过通货膨胀调整的零售额数据。开始，我们用 1972 年 2 月到 2000 年 12 月的实际零售额年化月度增长率的观测值来估计 AR（1）模型。我们估计如下等式：销售增长率_{*t*} = *b*₀ + *b*₁ 销售增长率_{*t-1*} + ϵ_t 。表 10-15 给出了该模型的估计结果。

表 10-15 月度实际零售额增长率：AR（1）模型（1972 年 2 月~2000 年 12 月）

回归统计量			
R 平方		0.065 9	
标准误		2.352 0	
观测数		347	
杜宾—沃特森		2.100 8	
	系数	标准误	<i>t</i> 统计量
截距	1.217 0	0.135 7	8.966 6
滞后 1 期	-0.257 7	0.052 2	-4.934 9
残差的自相关系数			
滞后期限	自相关系数	标准误	<i>t</i> 统计量
1	-0.055 2	0.053 7	-1.028 8
2	-0.153 6	0.053 7	-2.861 9
3	0.177 4	0.053 7	3.304 4
4	-0.102 0	0.053 7	-1.899 6
5	-0.132 0	0.053 7	-2.458 2
6	-0.267 6	0.053 7	-4.984 1
7	-0.136 6	0.053 7	-2.545 5
8	-0.092 3	0.053 7	-1.718 6
9	0.165 5	0.053 7	3.083 2
10	-0.173 2	0.053 7	-3.227 3
11	-0.062 3	0.053 7	-1.159 7
12	0.873 9	0.053 7	16.279 3

资料来源：U. S. Department of Commerce.

该模型残差的自相关系数在表 10-15 的底部给出，其结果表明季节性效应在该模型中非常显著。在 347 个观测值和 2 个参数的情况下，这个模型具有 345 个自由度。在 0.05 的显著性水平下，*t* 统计量的临界值约为 1.97。第 12 阶滞后的自相关系数（季节性自相关系数，因为我们使用月度数据）的值为 0.873 9，其 *t* 统计量为 16.28。该自相关系数的 *t* 统计量大于临界值（1.97），这说明我们应该拒绝第 12 阶自相关系数等于 0 的原假设。同样注意到许多其他的在表中所给出的自相关系数的 *t* 统计量都显著不同于 0。因此，表 10-15 所给出了的模型是设定错误的，所以我们无法依靠它来预测销售增长率。

假如我们在 AR（1）模型中增加销售增长率的季度滞后项（第 12 阶滞后项），再对于等式销售增长率_{*t*} = *b*₀ + *b*₁ 销售增长率_{*t-1*} + *b*₂ 销售增长率_{*t-12*} + ϵ_t 进行估计。表 10-16 给出了该等式估计的结果。季节性自相关系数的估计值（第 12 阶自相关系数）已经下降到了 0.033 5。前 12 个自相关系数 *t* 统计量的绝对值都不大于 0.05 显著性水平下的临界值 1.97。我们可以认为

该模型的残差不存在明显的序列相关性。因为我们有理由可以相信该模型是设定正确的，所以我们可以使用它来预测零售额增长率。注意表 10-16 中的 R^2 为 0.814 9，远大于表 10-15 中的 R^2 （由不带有季节性滞后项的模型所计算出的）。

表 10-16 月度实际零售额增长率：带有季度性滞后项的 AR（1）模型（1972 年 2 月~2000 年 12 月）

回归统计量			
R 平方	0.814 9		
标准误	1.048 7		
观测数	347		
杜宾—沃特森	2.430 1		
	系数	标准误	t 统计量
截距	0.151 6	0.066 9	2.265 2
滞后 1 期	-0.049 0	0.023 9	-2.048 4
滞后 12 期	0.885 7	0.023 7	37.302 8
残差的自相关系数			
滞后期限	自相关系数	标准误	t 统计量
1	-0.069 9	0.053 7	-1.301 9
2	-0.010 7	0.053 7	-0.198 5
3	0.094 6	0.053 7	1.763 0
4	-0.055 6	0.053 7	-1.035 5
5	-0.031 9	0.053 7	-0.593 6
6	0.028 9	0.053 7	0.538 6
7	-0.093 3	0.053 7	-1.738 2
8	0.006 2	0.053 7	0.115 0
9	-0.011 1	0.053 7	-0.205 9
10	-0.052 3	0.053 7	-0.974 2
11	0.037 7	0.053 7	0.702 6
12	0.033 5	0.053 7	0.623 1

资料来源：U. S. Department of Commerce.

我们如何解释该模型中的系数？为了预测当月的零售额增长率，我们需要知道上个月的零售额增长率以及 12 个月之前的零售额增长率。如果零售额在上个月和 12 个月前保持不变，那么表 10-16 中的模型预测零售额将在当月以大约 15% 的年增长率增长。而如果上个月零售额以 5% 的年增长率增长，而 12 个月之前的零售额以 10% 的年增长率增长，那么表 10-16 中的模型预测当月的零售额将以 $0.1516 - 0.0490 (0.05) + 0.8857 (0.10) = 0.2377$ 或者 23.8% 的年增长率增长。

10.8 自回归移动平均模型

至此，我们已经给出了自相关模型和移动平均模型作为对于时间序列建模的其他方法。我们

所考虑例子中的时间序列一般能用一个简单的自相关模型进行解释（带有或者不带有季节性滞后项）。^①然而，一些统计学家主张使用更加一般的模型，自回归移动平均（autoregressive moving-average, ARMA）模型。主张采用 ARMA 模型的人认为这些模型可能会更好地拟合数据，并提供比简单的自回归（AR）模型更好的预测值。然而，正如我们在本节之后所讨论的，对于这些模型的估计和使用存在着严格的限制条件。因为你可能会遇到 ARMA 模型，我们下面提供一个简要的介绍。

一个 ARMA 模型同时结合了因变量的自相关滞后项以及移动平均的误差项。对于 p 个自回归项和 q 个移动平均项的模型（记为 ARMA (p, q)），其等式为

$$x_{t+1} = b_0 + b_1 x_t + \cdots + b_p x_{t-p} + \epsilon_t + \theta_1 \epsilon_{t-1} + \cdots + \theta_q \epsilon_{t-q} \quad (10-16)$$

$$E(\epsilon_t) = 0, E(\epsilon_t^2) = \sigma^2, E(\epsilon_t \epsilon_s) = 0 \quad \text{对于 } t \neq s$$

式中， $b_1, b_2, \dots, b_p$ 是自回归参数，而 $\theta_1, \theta_2, \dots, \theta_q$ 是移动平均参数。

估计和使用 ARMA 模型具有许多限制条件。首先，ARMA 模型的参数可能非常不稳定。特别地，数据样本的微小变动或者对于 ARMA 参数的初始猜测可能导致非常不同的最终 ARMA 参数的估计值。其次，选择正确的 ARMA 模型不仅是一门科学，更是一门艺术。对于特定时间序列的确定 p 和 q 的标准不是非常完善。而且，即使在选择一个模型之后，该模型也可能预测不佳。

再重申一下，ARMA 可能非常不稳定，其依赖于所使用的数据样本和所估计的特定 ARMA 模型。因此，你应该对于那些声称某个 ARMA 模型比任何其他 ARMA 模型提供时间序列更好预测的说法持有怀疑的态度。事实上，在大部分情况下，你可以使用几乎不怎么复杂的 AR 模型来获得与那些 ARMA 模型相同准确程度的预测值结果。即使一些最坚定的 ARMA 模型倡导者们也认为这些模型不应该用于少于 80 个观测值的情况，并且他们也不建议使用 ARMA 模型来预测公司的季度销售额或者营业利润（即使使用 15 年的季度数据）。

10.9 自回归条件异方差模型

至今，我们忽视了时间序列模型中异方差的问题，并假设模型是同方差的。异方差性（Heteroskedasticity）是指误差项的方差依赖于自变量；同方差性（homoskedasticity）是指误差项的方差与自变量相互独立。我们之前一直假设误差项的方差是常数，它并不依赖于其自身时间序列的值或者先前误差项的大小。然而，有时该假设会被违背，而误差项的方差将不再是常数。在这种情况下，AR、MA 或者 ARMA 模型中回归系数的标准误将是不正确的，并且我们的假设检验也将是无效的。因此，我们可能会基于这些检验做出错误的投资决策。

例如，假设你对于一家公司的销售额构建了一个自回归模型。如果异方差存在，那么你模型回归系数的标准误将是错误的。由于异方差，一个或者多个销售额滞后项很可能表现出统计上的显著性，而实际上却并非如此。因此，如果你使用该模型来进行决策，那么你可能会得到一些次优的决策。

部分因为其工作而使其分享 2003 年诺贝尔经济学奖得主罗伯特·恩格尔（Robert F. Engle）在 1982 年首次提出了一种检验某种时间序列模型在某一期误差项的方差是否依赖于前一期误差项方差的方法。他将这种异方差称为自回归条件异方差（autoregressive conditional heteroskedasticity, ARCH）。

作为一个例子，考虑一个 ARCH (1) 模型

^① 对于标准普尔 500 的收益率（见例 10-14），我们选择了移动平均模型而不是自回归模型。

$$\epsilon_t \sim N(0, a_0 + a_1 \epsilon_{t-1}^2) \tag{10-17}$$

式中， ϵ_t 以其前一期的值为条件，它的分布是均值为 0，方差为 $a_0 + a_1 \epsilon_{t-1}^2$ 的正态分布。如果 $a_1 = 0$ ，那么每期误差项的方差仅为 a_0 。这时，方差在各期是常数，并且其不依赖于历史误差值。现在假设 $a_1 > 0$ 。那么某一期误差项的方差将依赖于前一期平方误差项有多大。如果某一期有一个较大的误差，那么在下一期误差项的方差将会更大。

恩格尔说明我们可以通过将先前从时间序列模型（AR，MA，或者 ARMA）所估计出的残差平方与一个常数和残差平方项的滞后项进行回归来检验一个时间序列是否是 ARCH（1）。我们可以估计线性回归方程

$$\hat{\epsilon}_t^2 = \alpha_0 + \alpha_1 \hat{\epsilon}_{t-1}^2 + u_t \tag{10-18}$$

式中， u_t 是一个误差项。如果 α_1 的估计值在统计显著不同于 0，那么我们认为该时间序列是 ARCH（1）。如果一个时间序列模型具有 ARCH（1）的误差项，那么在 $t + 1$ 时刻误差项的方差可以利用公式 $\hat{\sigma}_{t+1}^2 = a_0 + a_1 \hat{\epsilon}_t^2$ 在 t 时刻进行预测。

例 10-17 检验月度通货膨胀率中的 ARCH（1）

分析师莉塞特·米勒想要检验 CPI 通货膨胀月度数据是否包含自回归条件异方差。她可以利用时间序列模型所得到的残差来估计式（10-18）。正如在例 10-8 中所讨论的，如果她对 1971~2000 年的月度 CPI 通货膨胀率建模，那么她将得到结论：可以用来对于样本外通货膨胀率进行预测的最好自回归模型为 AR（1）模型。表 10-17 给出了该模型的误差项是否是 ARCH（1）的检验结果。

表 10-17 检验 AR（1）模型中年化月度 CPI 通货膨胀率残差的 ARCH（1）（1971 年 2 月~2000 年 12 月）

回归统计量			
R 平方	0.137 6		
标准误	26.329 3		
观测数	359		
杜宾—沃特森	1.912 6		
	系数	标准误	t 统计量
截距	7.295 8	1.505 0	4.847 8
滞后 1 期	0.368 7	0.048 8	7.548 3

资料来源：U. S. Bureau of Labor Statistics.

因为前一期残差平方项系数的 t 统计量大于 7.5，所以米勒容易拒绝误差项的方差不依赖于前一期误差项方差的原假设。因此，她在表 10-5 中所计算出的检验统计量是无效的额，她不应该使用它们来决定她的投资决策。

米勒的结论——月度通货膨胀率的 AR（1）模型在误差项中具有 ARCH 效应，这可能是由于所使用的样本区间的原因（1971~2000 年）。在例 10-9 中，她使用 1985~2000 年这个更短的样本区间，并得出月度 CPI 通货膨胀率服从 AR（1）过程。（这些结果在表 10-8 中给出。）表 10-17 表明对于整个样本区间（1971~2000 年）一个通货膨胀率的时间序列模型的误差项具有 ARCH 效应。那么更短样本区间（1985~2000 年）所估计出的误差项是否同样表现出 ARCH 效应呢？对于更短的样本区间，米勒利用月度通货膨胀率数据估计 AR（1）模型。[⊖]现在她要通过检验来观察误差项是否显示出 ARCH 效应。表 10-18 给出了其结果。

⊖ AR（1）的估计结果已经在例 10-9 中予以了给出。

表 10-18 检验 AR (1) 模型中年化月度 CPI 通货膨胀率残差的 ARCH (1) (1985 年 2 月 ~2000 年 12 月)

回归统计量			
R 平方		0.010 6	
标准误差		11.259 3	
观测数		191	
杜宾-沃特森		1.996 9	
	系数	标准误差	t 统计量
截距	5.393 9	0.922 4	5.847 9
滞后 1 期	0.102 8	0.072 4	1.420 5

资料来源: U. S. Bureau of Labor Statistics.

在这个样本中,前一期残差平方的系数很小,并只具有 1.420 5 的 t 统计量。因此,米勒无法拒绝该回归中的误差不具有自回归条件异方差的原假设。这是另一个证据说明对于 1985 ~ 2000 年的数据,AR (1) 是一个不错的拟合模型。误差项的方差表现出同方差性,而米勒可以依赖该 t 统计量。该结果再一次证实了对于 1971 ~ 2000 年整个区间,单个 AR 过程是设定错误的(它没有很好地描述该数据)。

假设一个包含 ARCH (1) 误差项的模型。该事实的后果是什么呢?第一,如果 ARCH 存在,那么回归参数的标准误就将是不正确的。在存在 ARCH 效应的情况下,我们将需要使用广义最小二乘法[⊖]或者其他修正异方差的方法来正确地估计时间序列模型中参数的标准误。第二,如果 ARCH 效应存在,并且我们对其进行了建模,例如以 ARCH (1) 建模,那么我们可以预测误差项的方差。例如,假设我们想要利用表 10-17 中的估计参数来预测通货膨胀中误差项的方差: $\hat{\sigma}_t^2 = 7.295 8 + 0.368 7 \hat{\varepsilon}_{t-1}^2$ 。如果某一期的误差为 0%,那么下一期所预测的误差的方差为 $7.295 8 + 0.368 7 (0) = 7.295 8$ 。如果某一期的误差为 1%,那么下一期所预测的误差的方差为 $7.295 8 + 0.368 7 (1^2) = 7.664 5$ 。

恩格尔和其他研究者已经提出许多比 ARCH (1) 模型更一般化的模型,包括 ARCH (p) 以及广义自回归条件异方差 (generalized autoregressive conditional heteroskedasticity, GARCH) 模型。在一个 ARCH (p) 模型中,当期误差项的方差线性依赖于前 p 期的误差平方项: $\alpha_t^2 = a_0 + a_1 \varepsilon_{t-1}^2 + a_2 \varepsilon_{t-1}^2 + a_{t-1}^2 + \cdots + a_p \varepsilon_{t-p}^2$ 。GARCH 模型类似于对时间序列中的误差方差项的 ARMA 模型。正像 ARMA 模型一样,GARCH 模型可能过分复杂而且不稳定:它们的结果可能极大依赖于样本区间的选择和对于 GARCH 模型中参数初始猜测。采用 GARCH 模型的金融分析师应该十分注意这些模型可能会产生错误的程度,并且他们应该检验 GARCH 的估计量是否对于样本以及参数初始猜测的变动保持稳健。[⊖]

10.10 两个以上时间序列的回归

至今,我们只讨论了关于单个时间序列的时间序列模型。虽然在相关性和回归一章以及多元回归一章中我们使用了线性回归模型来分析不同时间序列间的关系,但是在那些章节中我们完全

⊖ 参见 Greene (2003)。
⊖ 对于更多关于 ARCH, GARCH, 和其他时间序列方差的模型,参见 Hamilton (1994)。

忽视了单位根。一个包含单位根的时间序列不是协方差平稳的。如果任何线性回归中的时间序列包含一个单位根,那么回归中普通最小二乘估计值的检验统计量就可能是无效的。

为了确定我们是否能够对于多于一个时间序列用线性回归进行建模,让我们从单个自变量开始讨论;即有两个时间序列,一个对应于因变量,而另一个对应于自变量。于是我们将扩展我们的讨论到多个自变量。

我们首先使用单位根检验,例如迪克-福勒检验,对于这两个时间序列的每个进行检验,来确定它们是否具有一个单位根。^①对于这些检验的结果可能有几种可能的情况。一种可能的情况是我们发现这两个时间序列都没有单位根。于是我们可以安全地使用线性回归来检验这两个时间序列间的关系。不然,我们就可能不得不采用其他的检验,正如我们在本节后面将要讨论的。

例 10-18 单位根和费雪效应

在多元回归一章的例 9-8 中,我们通过估计通货膨胀率和美国无风险国库券收益率之间的回归关系来检验费雪效应。我们使用 1968 年第四季度到 2002 年第四季度的 137 个预期通货膨胀率和无风险国库券收益率的季度观测值的样本。我们使用线性回归来分析这两个时间序列间的关系。如果时间序列是协方差平稳的,那么该回归的结果将是有效的;即这两个时间序列都不具有单位根。因此,如果我们计算迪克-福勒的关于每个时间序列分别具有单位根的假设的 t 检验统计量,并发现我们可以拒绝无风险国库券收益率序列具有一个单位根的原假设和预期通货膨胀率时间序列具有单位根的原假设,那么我们就可以使用线性回归来分析这两个时间序列间的关系。在那种情况下,我们分析费雪效应的结果将是有效的。

第二种可能的情况是我们拒绝了自变量具有单位根的原假设而无法拒绝因变量具有单位根的原假设。在这种情况下,回归中的误差项将不是协方差平稳的。因此,一个或者多个下面的线性回归的假设将被违背:(1) 误差项的预期值为 0;(2) 所有观测值误差项的方差是常数;(3) 观测值间误差项是互不相关的。因此,所估计的回归系数和标准误将是不相合的。回归系数可能表现出显著,但是这些结果将是虚假的。^②因此,在这种情况下,我们不应该使用线性回归来分析这两个时间序列间的关系。

第三种可能的情况与第二种情况正好相反:我们拒绝了因变量具有单位根的原假设而无法拒绝自变量具有单位根的原假设。在这种情况下,与第二种情况一样,回归中的误差项也不是协方差平稳的,而我们也不能使用线性回归来分析这两个时间序列间的关系。

例 10-19 单位根和利用市盈率来预测股票市场收益率

约翰·德·弗里斯 (Johann de Vries) 正在分析南非股票市场的表现。他检验约翰内斯堡股票交易所 (Johannesburg Stock Exchange, JSE) 全股票指数的百分比变动是否可以用市盈率 (股价收益比, price-to-earnings ratio) (P/E) 来预测。使用 1983 年 1 月到 2002 年 12 月的月度数据,他使用 $(P_t - P_{t-1})/P_{t-1}$ 作为因变量,而用 P_{t-1}/E_{t-2} 作为自变量来进行回归,其中 P_t 是 JSE 指数在 t 时刻的数值,而 E_t 是该指数的收益。

① 对于单位根检验的理论细节,参见 Greene (2003) 或者 Hamilton (1994)。单位根检验可以在一些计量软件包,如 Eviews 中进行。

② 不平稳时间序列的虚假回归问题最早由 Granger 和 Newbold (1974) 进行了讨论。

德·弗里斯发现回归系数在统计上是显著的，而该回归的 R 平方的值很高。那么他在接受该回归是有效的之前应该采用什么其他的分析呢？

德·弗里斯需要对于这两个时间序列的每一个进行单位根检验。如果这两个时间序列中的一个具有单位根，表明其是不平稳的，那么线性回归的结果将是没有意义的，而且无法使用该结果来得出 P/E 能够预测股票市场收益率的结论。^①

下面一种可能的情况是两个时间序列都具有一个单位根。在这种情况下，我们需要在使用回归分析之前确定这两个时间序列是否是协整（cointegrated）的。^②如果一个长期的金融或者经济关系在两个时间序列间存在并使得它们在长期中不会无限偏离其中的某一个，那么这两个时间序列就是协整的。例如，如果两个时间序列具有一个共同的趋势，那么它们就是协整的。

在第四种情况中，两个时间序列都具有一个单位根，但是它们不是协整的。在这种情况下，正如上面的第二种和第三种情况，线性回归中的误差项将不是协方差平稳的，一些回归的假设将被违背，回归系数和标准误将不是相合的，故我们不能使用它们来进行假设检验。因此，一个变量关于另一个变量的线性回归将是毫无意义的。

最后，第五种可能的情况是两个时间序列都具有一个单位根，但是它们是协整的。在这种情况下，一个时间序列关于另一个时间序列的线性回归的误差项将是协方差平稳的。相应地，回归系数和标准误将是相合的，而我们也可以使用它们来做假设检验。然而，我们在解释协整变量的回归结果时应该非常小心。协整回归估计了两个序列长期的关系，但可能不是最好的估计这两个序列短期关系的模型。协整序列的短期模型（误差修正模型）在恩格尔（Engle）和格兰杰（Granger）（1987）和汉密尔顿（Hamilton）（1994）中进行了讨论，但是这些是专业的问题了。

现在让我们看一下如何来检验两个具有一个单位根的时间序列（正如上面最后两种情况所述的）之间的协整关系。^③恩格尔和格兰杰建议采用如下检验方法：如果 y_t 和 x_t 都是具有一个单位根的时间序列，我们应该操作如下。

1. 估计回归 $y_t = b_0 + b_1 x_t + \varepsilon_t$ 。
2. 利用迪克-福勒检验来检验第一步回归中的残差项是否具有一个单位根。因为残差基于的是回归的估计系数，所以我们无法使用标准的迪克-福勒检验的临界值。而我们必须使用由恩格尔和格兰杰计算出的临界值，该临界值考虑了关于迪克-福勒检验分布的回归参数的不确定性影响。
3. 如果（恩格尔-格兰杰）迪克-福勒检验无法拒绝误差项具有一个单位根的原假设，那么我们认为回归中的误差项不是协方差平稳的。因此，这两个时间序列不是协整的。在这种情况下，任何关于这两个序列的回归关系是虚假的。
4. 如果（恩格尔-格兰杰）迪克-福勒检验能够拒绝误差项具有一个单位根的原假设，那么我们认为回归中的误差项是协方差平稳的。因此，这两个时间序列是协整的。线性回归的参数和标准误将是相合的，而且使我们可以对于这两个序列的长期关系进行假设检验。

① Barr 和 Kantor（1999）包含了 P/E 时间序列是不平稳的证据。

② Engle 和 Granger（1987）首先提出了协整的概念。

③ 考虑一个具有一个单位根的时间序列 x_t 。对于许多这种金融和经济上的时间序列，该序列的一阶差分 $x_t - x_{t-1}$ 是平稳的。我们认为该序列，其一阶差分是平稳的，具有单个单位根。然而，对于一些时间序列，即使一阶差分也可能是非平稳的，它们需要进一步的差分才可能实现平稳。这类时间序列被认为具有多个单位根。在本节中，我们只讨论每个非平稳序列具有一个单位根的情况（而这一情况也是非常普遍的）。

例 10-20 检验英特尔公司销售额和名义 GDP 之间的协整关系

假如我们想要检验英特尔公司销售额的自然对数和 GDP 的自然对数之间是否是协整的 (即 GDP 与英特尔公司销售额之间是否存在一个长期的关系)。我们想要利用 1985 年第一季度到 1999 年第四季度的季度数据来检验该假设。下面是其具体步骤:

1. 检验这两个序列中的每一个是否具有一个单位根。如果我们无法拒绝两个序列都具有一个单位根的原假设,其表明两个序列都不是平稳的,那么我们必须接着检验这两个序列是否是协整的。

2. 在确定每个序列都具有一个单位根之后,我们估计回归 $\ln(\text{英特尔公司销售额}_t) = b_0 + b_1 \ln \text{GDP}_t + \varepsilon_t$,接着使用估计回归的残差进行关于该回归的误差项具有一个单位根的 (恩格尔—格兰杰) 迪克—福勒检验。如果我们拒绝回归的误差项具有一个单位根的原假设,那么我们拒绝不存在协整关系的原假设。即这两个序列是协整的。如果这两个时间序列是协整的,那么我们可以使用线性回归来估计英特尔公司销售额的自然对数和 GDP 的自然对数间的长期关系。

我们目前已经讨论了单个自变量的模型。我们现在将我们的讨论扩展到具有两个或者更多自变量的模型,此时我们会具有三个或者更多的时间序列。最简单的可能性是模型中没有一个时间序列具有单位根。于是,我们可以放心地使用多元回归来检验这些时间序列间的关系。

例 10-21 单位根和富达精选科技股股票基金的收益率

在多元回归的例 9-3 中,我们使用 1998 年 1 月到 2002 年 12 月的 60 个月度观测值的多元线性回归对于标准普尔 500/BARRA 成长型股票指数收益率或者标准普尔 500/BARRA 价值型股票指数收益率是否能解释富达精选科技股股票基金收益率进行了检验。当然,如果这三个时间序列中的任何一个具有一个单位根,那么我们回归分析的结果就可能是无效的。因此,我们可以使用一个迪克—福勒检验来确定是否这些序列中的某一个具有一个单位根。

如果我们拒绝所有这三个序列具有单位根的原假设,那么我们可以使用线性回归来分析这些序列间的关系。在这种情况下,我们对于影响富达精选科技股股票基金的收益率的因素的分析结果将是有效的。

如果至少一个时间序列 (因变量或者自变量之一) 具有一个单位根,而至少一个时间序列 (因变量或者自变量之一) 不具有单位根,那么回归中的误差项不可能是协方差平稳的。因此,在这种情况下,我们不应该使用多元线性回归来分析时间序列间的关系。

另一种可能情况是每个时间序列,包括因变量和每个自变量,都具有一个单位根。如果这种情况发生,那么我们需要确定这些时间序列是否是协整的。为了检验协整,我们采用与单个自变量模型类似的步骤。首先估计回归 $y_t = b_0 + b_1 x_{1t} + b_2 x_{2t} + \cdots + b_k x_{kt} + \varepsilon_t$ 。接着利用回归估计的残差进行关于回归误差项具有一个单位根原假设的 (恩格尔—格兰杰) 迪克—福勒检验。

如果我们无法拒绝回归的误差项具有一个单位根的原假设,那么我们无法拒绝不存在协整的原假设。在这种情况下多元回归的误差项不是协方差平稳的,故我们无法使用多元回归来分析时间序列间的关系。

如果我们能拒绝回归的误差项具有一个单位根的原假设,那么我们能拒绝不存在协整的原假设。然而,对于三个或者更多的协整的时间序列建模可能是困难的。例如,一个分析师可能想要基于国家的 GDP 和超过 65 岁的总人口数来预测一个退休服务公司的销售额。虽然该公司销售额、

GDP 以及超过 65 岁的人口数的每个序列具有一个单位根，而且是协整的，但是对于这三个序列的协整建模可能是困难的，而且这么做也超出了本书的讨论范围。没有掌握所有这些复杂问题的分析师应该避免对于多个具有单位根的时间序列建模进行预测：它们的回归系数可能是不相合的，而且可能产生不正确的预测值。

10.11 时间序列的其他议题

时间序列分析是一个广泛的话题，而且包含许多非常复杂的问题。本章我们的目的是提出这些对于金融分析师来说是最重要并且是相对容易处理的时间序列问题。在本节中，我们简单讨论一些我们之前没有讨论但是可能对于分析师非常有用的问题。

在本章中，我们已经给出了如何使用时间序列模型来进行预测。我们也介绍了 RMSE 作为比较预测模型的一个准则。然而，我们还没有讨论如何来度量使用时间序列模型所做出预测值的不确定性。这些预测值的不确定性可能非常大，所以我们应该在做出投资决策时加以考虑。幸运的是，我们可以使用评价线性回归模型预测值不确定性的相同的方法来评价时间序列预测值的不确定性。为了准确评价预测的不确定性，我们需要考虑时间序列模型中误差项的不确定性和估计参数的不确定性。当使用包含多个自变量的回归时，评价这种不确定性将变得相当复杂。

在本章中，我们使用美国 CPI 通货膨胀率序列来介绍分析师们在实际使用时间序列模型中所面临的一些挑战。我们使用美联储政策的信息探讨了将通货膨胀率序列分割成两个时间段所造成的后果。在处理金融时间序列的工作中，我们可能怀疑一个时间序列具有多种状态，但是没有信息能试图将数据分成不同的状态。如果你面临这类问题，那么你可能想要对该单个时间序列本身利用其他方法进行检验，特别是状态切换回归模型，以识别出多种状态。

如果你对于这些或者其他更高级的时间序列的问题感兴趣，你可以在戴博德（Diebold）（2004）和汉密尔顿（Hamilton）（1994）中学到更多内容。

10.12 时间序列预测建议采取的步骤

下面给出的是建立一个预测时间序列的模型所应该采取的具体步骤。

1. 理解你所面临的投资问题，并做出初始的模型选择。一种选择是回归模型，该模型能基于与其他变量间假设存在的因果关系来预测一个变量的未来变化情况。另一种选择是时间序列模型，该模型试图基于同一个变量历史的变化情况来预测该变量未来的变化。

2. 如果你已经决定使用一个时间序列模型，那么你要整理该时间序列，并将其绘制在图上以观察它是否看上去是协方差平稳的。该图可能反映出该序列相对于协方差平稳的重大偏离，具体表现为下述情况：

- 一个线性趋势；
- 一个指数趋势；
- 季节性效应；
- 时间序列在样本区间中的一个显著变动（例如，均值或者方差的一个变化）。

3. 如果你发现时间序列中没有显著的季节性效应或者变动，那么很可能一个线性趋势或者一个指数趋势将可以充分对于该时间序列建模。在这种情况下，你应该采取下列步骤：

- 确定一个线性或者指数趋势是否看上去非常合理（一般是通过对该序列作图来进行判断）。
- 估计该趋势。
- 计算其残差。
- 使用杜宾-沃特森统计量来确定该残差是否具有显著的序列相关性。如果你发现残差没有显著的序列相关性，那么趋势模型就足以刻画时间序列的动态变化，而你也可以使用该模型进行预测。

4. 如果你发现趋势模型的残差中存在显著的序列相关，那么就应该使用一个更复杂的模型，如一个自回归模型。然而，首先你要再次检验时间序列是否是协方差平稳的。下面列出的是一组平稳性被违背的情况，同时我们也给出了潜在的调整这些时间序列以使其协方差平稳的处理方法：

- 如果该时间序列具有一个线性趋势，则对其进行一阶差分。
- 如果该时间序列具有一个指数趋势，则对其取自然对数后进行一阶差分。
- 如果该时间序列在样本区间有显著的变动，则对于该变动之前和之后的不同的时间序列模型进行估计。
- 如果该时间序列具有明显的季节性效应，则在模型中包含季节性滞后项（在下面第7步中汇总进行讨论）。

5. 在你成功地将一个原始的时间序列转化为一个协方差平稳的时间序列之后，你一般可以使用一个较短的（阶数较小的）自相关模型对于转化后的序列进行建模。[⊖] 为了确定使用哪一个自回归模型，采取下列步骤：

- 估计一个 AR (1) 模型。
 - 检验该模型的残差是否具有明显的序列相关性。
 - 如果你发现该残差不具有序列相关性，那么你可以使用 AR (1) 模型进行预测。
6. 如果你发现残差中显著的序列相关性，那么，接着使用 AR (2) 模型并检验 AR (2) 模型残差的序列相关性是否显著。
- 如果你发现不具有显著的序列相关性，那么使用 AR (2) 模型。
 - 如果你发现残差中显著的序列相关性，那么继续增加 AR 模型的阶数直至残差的序列相关性不再显著。
7. 你的下一步是检查季节性效应。你可以使用下面两种方法之一：
- 对数据作图，并检查是否存在有规则的季节性模式。
 - 检查数据以了解 AR 模型残差的季节性自相关系数是否显著（例如，季度数据的第四阶自相关系数）以及季节性自相关系数之前和之后的自相关系数是否显著。为了修正季节性效应，你可以在你的 AR 模型中加入季节性滞后项。例如，如果你使用季度数据，你可以在一个 AR (1) 或者 AR (2) 模型中加入时间序列的四阶滞后项作为另外一个变量。
8. 接下去，检验残差是否具有自回归条件异方差。例如，为了检验 ARCH (1)，可以进行以下操作：

⊖ 大部分金融时间序列可以用自回归过程进行建模。对于一些时间序列，一个移动平均模型可能拟合得更好。为了了解为什么是这样的，你可以检验时间序列的前五阶或者前六阶的自相关系数。如果自相关系数在 q 阶之后突然降至 0，那么一个 (q 阶的) 移动平均模型是合适的。如果一个自相关系数开始很大，而之后缓慢下降，那么一个自回归模型是合适的。

- 将你时间序列模型所得到的残差平方与其残差平方的滞后项进行回归。
- 检验残差平方滞后项前的系数是否显著不为 0。
- 如果残差平方滞后项前的系数不是显著不为 0，那么残差就没有表现出 ARCH 效应，而你也可以使用你时间序列估计的标准误。
- 如果残差平方滞后项前的系数显著不为 0，那么你应该使用广义最小二乘法或者其他方法来修正 ARCH 效应。

9. 最后，你也可能想要对于模型样本外预测的表现进行检验，以了解该模型相对于其样本内预测的表现，其样本外预测表现是如何的。

依次进行这些步骤，你就可以有理由确信你的模型是设定正确的。

投资组合的概念

11.1 引言

在定量投资分析的各项研究中，没有一项研究内容会像投资组合理论那样被人们广泛地考察和热烈地探讨。投资组合经理在近 50 年来已经研究过的问题包括：

- 投资组合的哪些特征是重要的，我们应如何来量化它们？
- 我们应该如何对风险进行建模？
- 如果我们能够了解资产收益率的分布，那么我们将如何来选择一个最优的投资组合？
- 一个投资组合中风险资产和无风险资产最优的组合方式应该是怎样的？
- 使用历史收益率数据来预测一个投资组合未来特征的局限性是什么？
- 除了市场风险之外，我们还应该考虑哪些风险因素？

在本章中，我们将介绍一些有助于投资组合管理的重要定量方法。在第 2 节中，我们会关注均值一方差分析以及相关的模型和问题。接着在第 3 节中，我们会利用均值一方差分析来回答一些我们所遇到的问题，并给出应对这些问题的方法。我们将介绍单因素模型，该模型是一个市场模型，它将用市场指数这个变量来解释资产的收益率。在第 4 节中，我们会给出一个利用多个因素来解释资产收益率的模型，同时，我们还会列举一些在目前实务中对于这些模型的重要应用。

11.2 均值方差分析

投资组合分散化在什么情况下才会降低风险？是否存在所有风险厌恶的投资者都会避而远之的投资组合？这些问题都是哈利·马克维茨（Harry Markowitz）在他的研究中所讨论的问题，而该研究也使其分享了 1990 年的诺贝尔经济学奖。

作为历史最悠久、可能也是最为人们所接受的现代投资组合理论中的部分，均值一方差投资组合理论为通过考察风险和收益来进行投资组合选择的行为提供了理论基础。在本节中，我们将对马克维茨理论进行讨论。我们会通过几个例子来给出投资组合分散化的一些原则，并讨论该理

论在实际应用中所存在的几个重要问题。

均值一方差投资组合理论基于的是这样的一个想法：投资机会的价值可以用收益率的均值和方差来进行度量。马克维茨将这一建立投资组合的方法称为均值一方差分析（mean-variance analysis）。均值一方差分析基于如下假设：

- 1. 所有投资者都是风险厌恶的；他们在相同的期望收益率的情况下，更愿意承担较小的风险。[Ⓐ]
- 2. 所有资产的期望收益率都是已知的。
- 3. 所有资产收益率的方差和协方差都是已知的。
- 4. 投资者只需要知道期望收益率、收益率的方差以及协方差就可以确定出最优的投资组合。他们会忽视收益率的偏度、峰度以及其他的分布特征。[Ⓑ]
- 5. 不存在交易成本和税收。

注意第 1 个假设并不意味着所有投资者都具有相同的风险容忍程度。投资者愿意接受的风险水平是不同的；但是，风险厌恶的投资者在给定的期望收益率水平下，都更愿意承担较小的风险。在现实中，期望收益率以及资产收益率的方差和协方差都是未知的，但是我们可以对它们进行估计。当我们使用均值一方差分析时，对这些量估计的偏差可能是导致决策中出现错误的一个原因。

第 4 个假设是非常重要的一个假设，因为它认为我们可能只会依赖于资产收益率分布中某些特定的描述性统计量——期望收益率，方差以及协方差，来决定投资组合的最优组合方式。

11.2.1 最小方差前沿及其相关概念

投资者使用均值一方差方法进行投资组合选择的目的是选择一个最有效的投资组合。一个有效投资组合（efficient portfolio）是指一个在给定的（以收益率的方差或者标准差度量的）风险水平下提供最高期望收益率的投资组合。因此，如果一个投资者使用标准差来度量其对于风险的容忍程度，那么她会选择这样的—一个投资组合，该组合收益率的标准差与其风险容忍程度是一致的，而其期望收益率（根据她估计的结果）是在该风险水平下最高的。下面，我们首先从 2 种资产类（政府债券和大盘股）中构造投资组合的情形来开始对于投资组合选择问题的讨论。

表 11-1 给出了我们对于这 2 种资产期望收益率以及其收益率的标准差和它们收益率之间相关系数的假定值。

表 11-1 假定的期望收益率、方差和相关系数：2 种资产的情形

	资产 1 大盘股	资产 2 政府债券		资产 1 大盘股	资产 2 政府债券
期望收益率	15%	5%	标准差	15%	10%
方差	225	100	相关系数	0.5	

在对于有效投资组合的求解过程中，我们首先必须找出在每个给定期望收益率水平下具有最小方差的投资组合。这类投资组合被称为最小方差投资组合（minimum-variance portfolios）。正如我们将会看到的，有效投资组合的集合是最小方差投资组合集合的一个子集。

Ⓐ 更多关于风险厌恶的概念以及该概念在投资组合中所发挥的作用，可以参见如 Sharpe, Alexander 和 Bailey (1999) 或者 Reilly 和 Brown (2003)。
Ⓑ 这一假设可以通过设定收益率服从一个正态分布或者假定投资者拥有一个在数学上完全可以由均值和方差来表示的收益率和风险的偏好来得到。

我们从表 11-1 中可以看到大盘股（资产 1）收益率的标准差为 15%，政府债券（资产 2）收益率的标准差为 10%，而这 2 种资产收益率间的相关系数为 0.5。因此，我们可以计算出关于投资于大盘股比例（ ω_1 ）以及投资于政府债券比例（ ω_2 ）的投资组合收益率的方差。因为投资组合只包含这 2 种资产，所以我们有关系式 $\omega_1 + \omega_2 = 1$ 。当投资组合 100% 投资到资产 1 中时， ω_1 为 1.0 而 ω_2 为 0；而当 ω_2 为 1.0 时，那么 ω_1 就为 0，此时，投资组合就被 100% 投资到资产 2 之中。另外，当 ω_1 为 1.0 时，我们就知道投资组合的期望收益率和方差就是资产 1 的期望收益率和方差。相反，如果 ω_2 为 1.0，那么，投资组合的期望收益率和方差就是资产 2 的期望收益率和方差。在这种情况下，如果投资组合 100% 投资于大盘股，那么投资组合达到最大的期望收益率，其值为 15%；而如果投资组合 100% 投资于政府债券，那么投资组合达到最小的期望收益率，其值为 5%。

在我们能够确定出包含大盘股和政府债券的所有投资组合的风险和收益之前，我们必须要知道任意包含 2 个资产的投资组合的期望收益率，方差和标准差是如何依赖于这 2 种资产的期望收益率、收益率的方差和它们收益率之间相关系数的。

对于任一个包含 2 种资产的投资组合，其期望收益率 $E(R_p)$ 为

$$E(R_p) = \omega_1 E(R_1) + \omega_2 E(R_2)$$

式中， $E(R_1)$ 表示资产 1 的期望收益率

$E(R_2)$ 表示资产 2 的期望收益率

该投资组合收益率的方差为

$$\sigma_p^2 = \omega_1^2 \sigma_1^2 + \omega_2^2 \sigma_2^2 + 2\omega_1 \omega_2 \rho_{1,2} \sigma_1 \sigma_2$$

式中， σ_1 表示资产 1 收益率的标准差

σ_2 表示资产 2 收益率的标准差

$\rho_{1,2}$ 表示资产 1 和资产 2 收益率之间的相关系数

而 $\text{Cov}(R_1, R_2) = \rho_{1,2} \sigma_1 \sigma_2$ 表示这 2 种资产收益率间的协方差，回忆一下，相关系数的定义，2 个变量的相关系数等于它们的协方差除以它们各自的标准差。投资组合收益率的标准差为

$$\sigma_p = (\omega_1^2 \sigma_1^2 + \omega_2^2 \sigma_2^2 + 2\omega_1 \omega_2 \rho_{1,2} \sigma_1 \sigma_2)^{1/2}$$

在这个例子中，投资组合的期望收益率为 $E(R_p) = \omega_1 (0.15) + \omega_2 (0.05)$ ，而投资组合收益率的方差为 $\sigma_p^2 = \omega_1^2 0.15^2 + \omega_2^2 0.10^2 + 2\omega_1 \omega_2 (0.5) (0.15) (0.10)$ 。

给定这两种资产的期望收益率、收益率方差以及它们的收益率间的相关系数，我们就可以确定出投资组合的期望收益率和方差，并将其表示为关于投资于大盘股和政府债券的资产比例的一个函数。表 11-2 给出了当大盘股的投资权重从 0 变动到 1.0 时的投资组合的期望收益率、方差以及标准差。

表 11-2 股票和债券投资组合的期望收益率和风险之间的关系

期望 收益率(%)	投资组合 的方差	投资组合的 标准差(%)	大盘股 比例(ω_1)	政府债券 比例(ω_2)	期望 收益率(%)	投资组合 的方差	投资组合的 标准差(%)	大盘股 比例(ω_1)	政府债券 比例(ω_2)
5	100.00	10.00	0	1.0	11	133.00	11.53	0.6	0.4
6	96.75	9.84	0.1	0.9	12	150.75	12.28	0.7	0.3
7	97.00	9.85	0.2	0.8	13	172.00	13.11	0.8	0.2
8	100.75	10.04	0.3	0.7	14	196.75	14.03	0.9	0.1
9	108.00	10.39	0.4	0.6	15	225.00	15.00	1.0	0
10	118.75	10.90	0.5	0.5					

正如表 11-2 所示，当大盘股的权重为 0.1 时，投资组合的期望收益率为 6%，收益率的方差为 96.75。[⊖]该投资组合比投资于股票的权重为 0 的投资组合（即完全投资于政府债券的投资组合）具有更高的期望收益率和更小的方差。这种在风险—收益上改善的特征揭示了分散化的威力：因为大盘股的收益率不是完全与政府债券相关的（它们的相关系数不为 1），所以通过在投资组合中加入一些大盘股，我们可以在增加期望收益率的同时降低收益率的方差，并且这种对于改善投资组合风险—收益特征的做法不会带来任何额外的成本。

图 11-1 绘制出了由政府债券和大盘股所组成的投资组合的所有可能的风险和收益率的情况。图 11-1 中的 y 轴表示投资组合的期望收益率，而 x 轴表示投资组合收益率的方差。

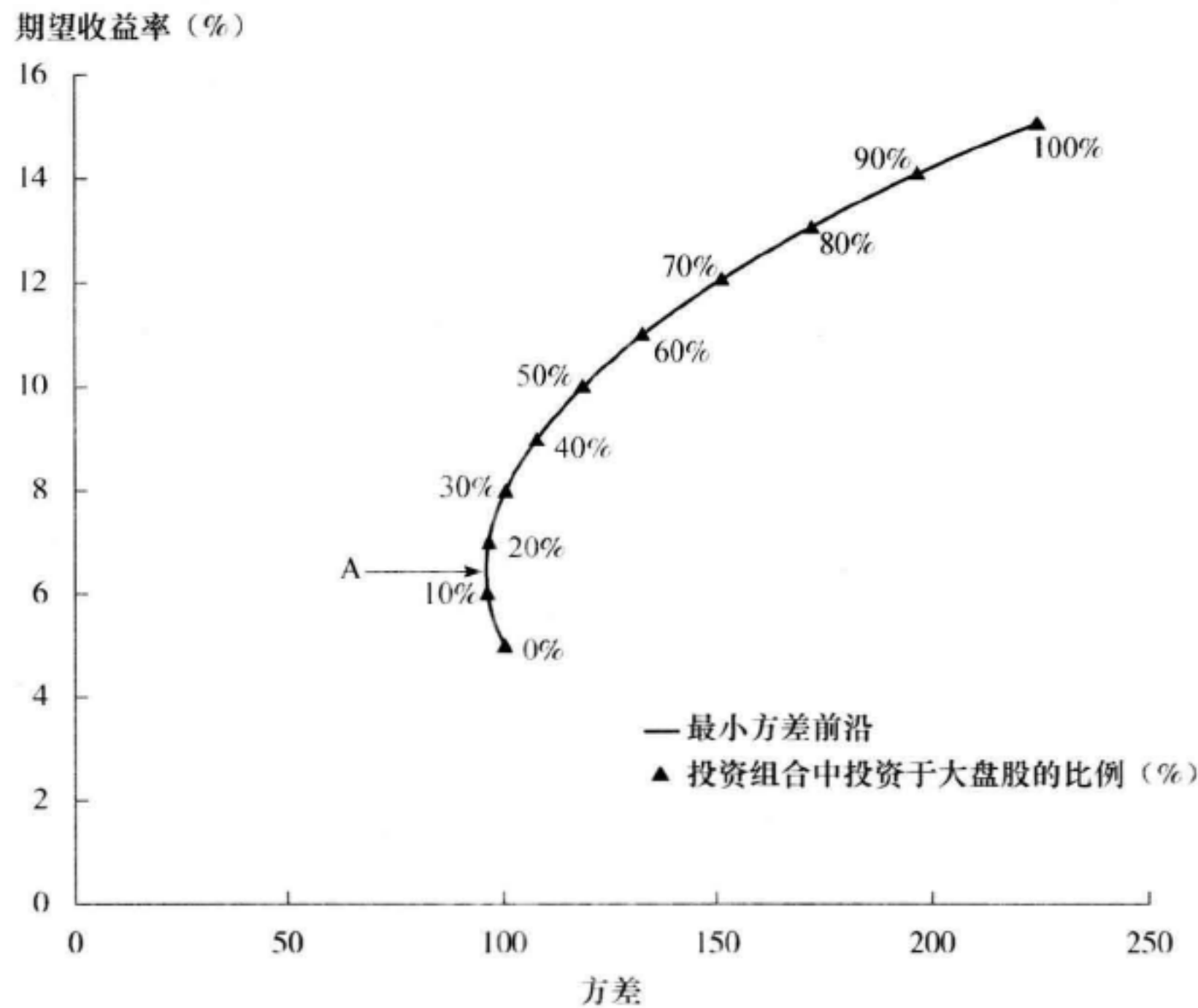


图 11-1 最小方差前沿：大盘股和政府债券

2 种资产进行组合的例子是非常特别的，因为所有包含这 2 种资产的投资组合都可以绘制在如图所示的一条曲线上（在给定期望收益率水平的情况下，曲线上存在唯一包含这 2 种资产的组合）。该曲线被称为投资组合可能曲线（portfolio possibilities curve）——一条由 2 种资产组合所形成的反映投资组合期望收益率和风险的曲线。我们也可以把图 11-1 中的这条曲线称为最小方差前沿（minimum-variance frontier），因为它反映了在一个给定的期望收益率水平下，投资组合可以达到的最小方差。最小方差前沿是一个比投资组合可能曲线更加有用的概念，因为它也被应用到包含多于 2 种资产的投资组合情况之中。在多于 2 种资产的更加一般的情形中，任何绘制在一条给定期望收益率水平的假想水平线上的投资组合都具有相同的期望收益率，而随着我们沿着这条水平线向左移动，我们会得到更小收益率方差的投资组合。而给定该期望收益率水平，在该水平线上最左侧可达到的投资组合就是最小方差投资组合，它是最小方差前沿上的一个点。在 3 种或者更多资产的情形中，最小方差前沿就是一个真正的前沿了：它是代表所有可能的期望收益率和风险组合区域的边界（可行

⊖ 注意 96.75 的单位是百分比的平方。用小数来表示，投资组合的期望收益率为 0.06，而方差为 0.009 675。

域的边界)。该区域是由这样的事实所产生的：在 3 种或者更多资产的情形中，有无限多个投资组合可以提供给定的期望收益率水平。^①在 3 种或者更多资产的情形下，如果我们从最小方差前沿上的一点向右移动，那么我们将达到另一个投资组合，而该组合具有更大的风险。

从图 11-1 中，我们注意到全局最小方差投资组合的方差（具有最小方差的一点）看上去接近于 96.43（点 A），此时投资组合的期望收益率为 6.43。该全局最小方差投资组合中 14.3% 的资产投资于大盘股，而 85.7% 的资产投资于政府债券。给定这些假定的收益率、标准差以及相关系数之后，一位投资组合经理不应该选择投资于大盘股比例小于 14.3% 的投资组合，因为任何这样的投资组合将比全局最小投资组合具有更高的方差和更小的期望收益率。所有在最小方差前沿上低于点 A 的点都劣于全局最小方差的投资组合，所以我们应该避免选择它们。

金融经济学家经常说位于全局最小方差投资组合（图 11-1 中的点 A）之下的投资组合被那些与其具有相同方差但是期望收益率更高的其他投资组合所占优。因为这些被占优的投资组合没有有效地利用风险，因此，它们是无效的投资组合。最小方差前沿中从全局最小方差组合开始向上延伸的部分被称为有效前沿（efficient frontier）。在有效前沿上的投资组合能够提供在其收益率方差水平上最大的期望收益率。有效投资组合充分地利用了风险：以收益率均值和方差进行投资组合选择的投资者将只会限制在有效前沿上进行投资组合的选择。这种方法使得我们所要考虑的投资组合数量大为减少，从而简化了投资组合选择的工作。如果一位投资者可以用收益率的方差和标准差来量化表示其风险容忍程度，那么在给定方差或者标准差水平上的有效投资组合将代表最优的均值一方差选择。

因为标准差比方差更容易解释，所以投资者经常会绘制期望收益率和标准差的图表，而不是期望收益率和方差的图表。^②图 11-2 将该例子中投资组合的期望收益率绘制在 y 轴上，而将其收益率的标准差标注在 x 轴上。^③所绘制出的曲线仍被称为最小方差前沿。

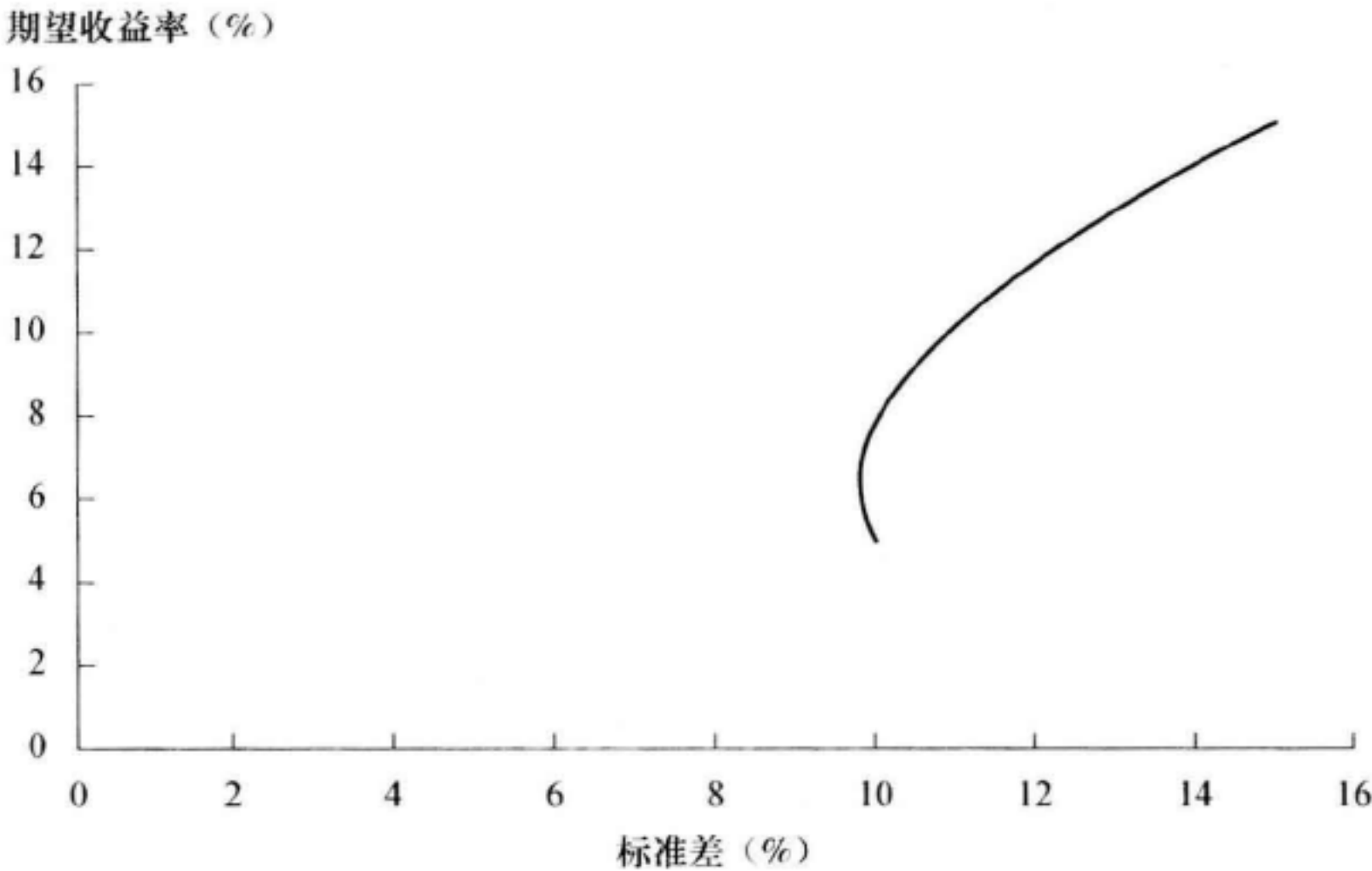


图 11-2 最小方差前沿：大盘股和政府债券

① 例如，如果我们具有 3 种资产，它们的期望收益率分别为 5%、12% 和 20%，而我们想要一个期望收益率为 11% 的投资组合，那么我们应该使用如下的等式来求解投资组合的权重（利用投资组合权重的和必须等于 1 的事实）： $11\% = (5\% \omega_1) + (12\% \omega_2) + [20\% (1 - \omega_1 - \omega_2)]$ 。这个含有 2 个未知量 ω_1 和 ω_2 的等式具有无穷多个可能的解，而每个解都代表一个可能的投资组合。

② 期望收益率和标准差是以相同的单位（百分比）来度量的。

③ 对于本章中剩下的部分，我们将只绘制期望收益率和收益率标准差的图像。

例 11-1 给出了确定一个历史的最小方差前沿的具体过程。

例 11-1 利用历史的美国收益率数据所得到的一个 2 种资产的最小方差前沿

苏珊·菲茨西蒙斯 (Susan Fitzsimmons) 已经决定将其退休计划的资产分别投资到一个美国小盘股的指数基金和一个美国长期政府债券的指数基金之中。菲茨西蒙斯决定使用均值一方差分析来帮助她确定投资到每种基金中资金的比例。假设利用 1970~2002 年的月度历史收益率能够准确地估计出期望收益率和方差,那么她可以计算出这 2 个准备投资的指数基金的平均收益率,收益率的方差以及收益率之间的相关系数。表 11-3 给出了这些历史统计量。

表 11-3 平均收益率和收益率的方差 (年化的, 基于月度数据, 1970 年 1 月~2002 年 12 月)

资产类别	平均收益率	方差
美国小盘股	14.63%	491.8
美国长期政府债券	9.55%	109.0
相关系数	0.138	

资料来源: Ibbotson Associates.

给定这些统计量,菲茨西蒙斯可以利用期望收益率和方差确定投资组合中在这 2 种资产间分配的比例。为了计算出这一结果,她必须计算出

- 投资组合可能的期望收益率的范围 (最小值和最大值);
- 每个可能的期望收益率水平下最小方差投资组合中这 2 种资产各自所占的比例;
- 每个可能的期望收益率水平下的方差[⊖]。

因为美国政府债券具有比美国小盘股更低的期望收益率,所以具有最小期望收益率的投资组合对美国长期政府债券的投资比例为 100%,而对美国小盘股的投资比例为 0%,其期望收益率为 9.55%。相反,最高期望收益率的投资组合对美国小盘股的投资比例为 100%,而对美国长期政府债券的投资比例为 0%,其期望收益率为 14.63%。因此,投资组合可能的期望收益率的范围是 9.55%~14.63%。

菲茨西蒙斯现在要确定在不同期望收益率水平 (从最小的期望收益率 9.55% 到最大期望收益率 14.63% 为止) 下投资到 2 种资产类别中的比例。在每个期望收益率水平下的权重确定了包含这 2 种资产类别投资组合的方差。表 11-4 给出了不同期望收益率水平下投资组合中资产的组成情况。

表 11-4 美国小盘股和美国长期政府债券的投资组合在最小方差前沿上的点

期望收益率 (%)	方差	标准差 (%)	小盘股权重 ω_1	政府债券权重 ω_2	期望收益率 (%)	方差	标准差 (%)	小盘股权重 ω_1	政府债券权重 ω_2
9.55	109.0	10.4	0.000	1.000	10.55	99.5	10.0	0.197	0.803
9.65	106.2	10.3	0.020	0.980	10.75	102.6	10.1	0.236	0.764
9.95	100.2	10.0	0.079	0.921	14.63	491.8	22.2	1.000	0.000
10.25	98.0	9.9	0.138	0.862					

⊖ 在 2 种资产的情形中,正如前面所述的,在给定期望收益率水平下存在着唯一提供该期望收益率的 2 种资产的组合,因此在给定期望收益率水平下也存在唯一投资组合的方差。所以每个期望收益率水平下所计算的投资组合方差就是该期望水平下最小方差投资组合的方差。

表 11-4 给出了我们从最小期望收益率变动到最大期望收益率过程中，投资组合中单个资产类别权重的变化情况。当期望收益率为 9.55% 时，长期政府债券的权重为 100%。随着我们增加期望收益率，长期政府债券的权重会逐渐下降；同时，美国小盘股的权重会上升。这一结果是有道理的，因为我们知道最大期望收益率 14.63% 必须对应于美国小盘股 100% 的权重。表 11-4 中的权重反映出了这个性质。注意到全局最小方差组合（也是全局最小标准差组合）同时包含了这 2 种资产。一个只包含债券的投资组合比全局最小方差组合具有更大的风险和更低的期望收益率，这是因为分散化能够降低总投资组合的风险，而我们在不久之后就要讨论这个问题。

图 11-3 通过将期望收益率作为标准差的函数，绘制出了 1970 ~ 2002 年这一区间的最小方差前沿（在 2 种资产类的例子中，该前沿只是一条投资组合的可能曲线）。

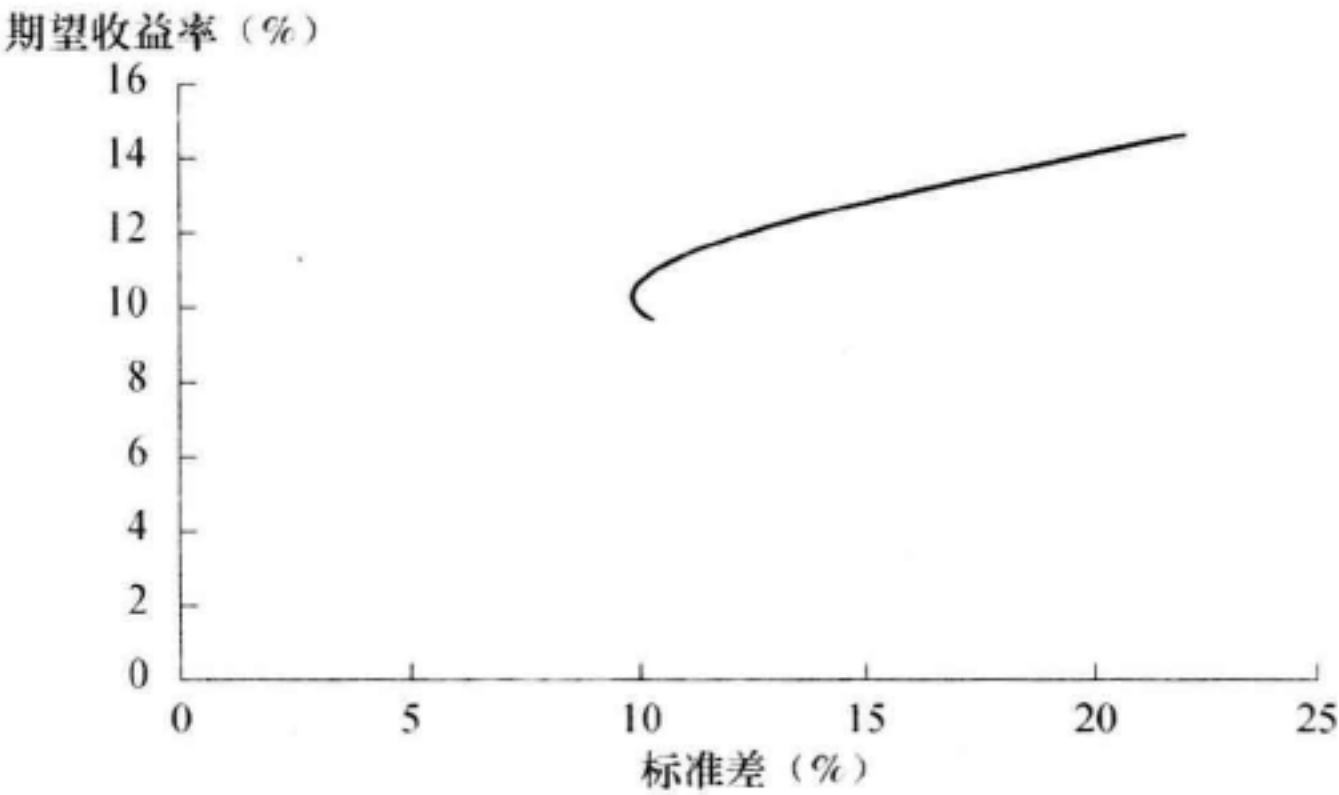


图 11-3 最小方差前沿：美国小盘股和政府债券

例如，如果菲茨西蒙斯用 10% 的标准差来反映其风险容忍程度，那么均值一方差分析建议她选择投资于小盘股的权重约为 0.20，而投资于长期政府债券的权重约为 0.80 的投资组合。我们在本章的后面将要讨论的一个重要注意事项是：即使是输入变量的很小变化都会对最小方差前沿产生较为显著的影响，而且未来情况可能与过去有明显的不同。历史记录只是用来建立计算最小方差前沿所需输入变量的一个起始点。[Ⓐ]

一个对于投资组合风险和收益率的权衡不仅依赖于资产的预期收益率和方差，还依赖于资产收益率间的相关系数。回到之前大盘股和政府债券的例子中，我们当时假设它们的相关系数为 0.5。而风险收益的权衡在其他相关系数值的情况下是不同的。图 11-4 给出了不同权重条件下的包含大盘股和政府债券的投资组合的最小方差前沿。[Ⓑ]在 4 种不同相关系数情况下，权重都会从投资于政府债券的 100% 和大盘股的 0% 变动到政府债券的 0% 和大盘股的 100%。在图 11-4 中给出的相关系数分别是 -1，0，0.5 和 1。

图 11-4 给出了许多最小方差前沿和分散化的有趣特征：[Ⓒ]

Ⓐ 同样注意到该历史数据是月度的，故其对应于一个月度的投资区间。如果我们使用一个不同的时间区间（比如季度区间），那么最小方差前沿就可能有很大的不同。

Ⓑ 回忆一下，在表 11-1 中，大盘股具有给定的期望收益率和标准差都为 15%，而政府债券具有给定的期望收益率为 5%，标准差为 10%。

Ⓒ 我们所考察的例子（或者我们的观测值）中的 2 种资产没有一个是相对另一个占优的。在均值一方差分析中，一个资产 A 被另一个资产 B 占优，当（1）资产 B 的平均收益率大于等于资产 A 的平均收益率，但是 B 的标准差小于 A 的标准差；或者（2）资产 B 的平均收益率严格大于资产 A 的平均收益率，但是 A 和 B 的收益率具有相同的标准差。连接 2 种互不占优资产的直线的斜率是正的。

- 所有前沿的端点都是相同的。这一事实不应该令人惊讶，因为在一个端点，所有的资产被投资于政府债券，而在另一个端点，所有的资产被投资到大盘股。在每个端点期望收益率和标准差都是相关资产（股票或者债券）的收益率和标准差。
- 当相关系数为 +1 时，最小方差前沿就是一条向上倾斜的直线。如果我们从该直线上的任一点开始，对于标准差每一个百分点的变动，我们都会获得相同增加量的期望收益率。在相关系数为 +1 的情况下，一个资产的收益率（不仅仅是期望收益率）是另一个资产收益率的严格为正的线性函数。[⊖]因为 2 种资产收益率的波动以这种方式相互影响，所以一个资产的收益率不能消除或者平滑另一种资产收益率的波动。因此，对于相关系数为 +1 的情况，分散化没有任何效果。

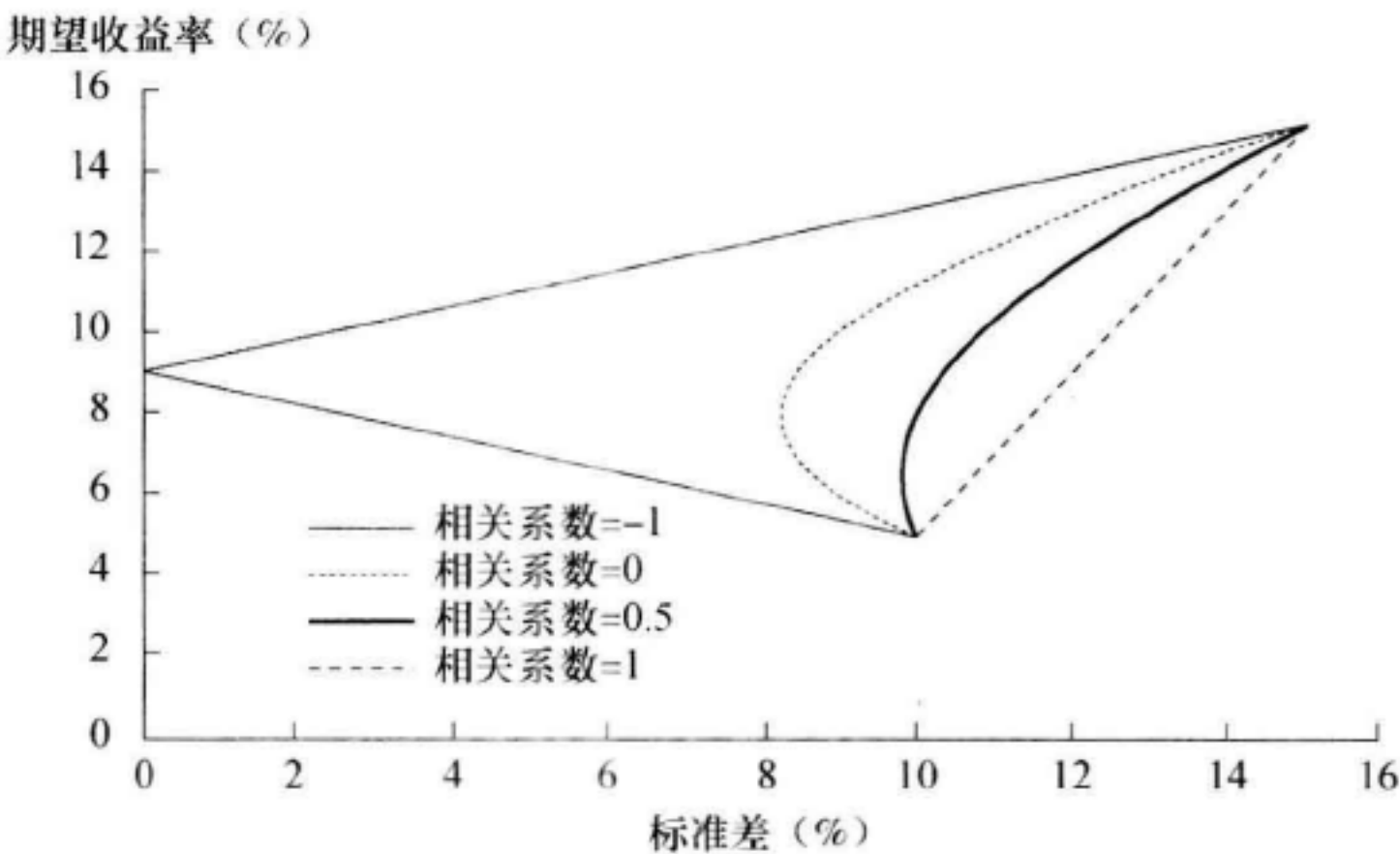


图 11-4 不同相关系数情况下的最小方差前沿：大盘股和政府债券

- 当我们将相关系数从 +1 的情况变动到 0.5 的情况时，最小方差前沿会向左侧（即标准差减少的方向）突出。对于相关系数小于 +1 的情况，我们可以在任何可以达到的期望收益率水平上取得比之前相关系数为 +1 的情况更小的标准差。随着我们从相关系数为 0.5 的情况变动到更小相关系数的情况，最小方差前沿将会更加严重地向左侧突出。
- 相关系数为 0.5，0 和 -1 情况下的前沿具有一段斜率为负的部分。[⊖]这意味着如果我们从最低点（100% 投资于政府债券）开始，逐渐将资金转移到股票之中，直到我们达到全局最小方差投资组合，那么我们能获得更大的期望收益率和较小的风险。因此，相对于最初完全投资于政府债券的头寸，这些相关系数情况的每一种都会获得分散化的好处。分散化的好处是指在不减少期望收益率的情况下通过分散化来降低投资组合收益率的标准差。因为最小方差前沿会随着我们降低相关系数而向左侧突出，所以我们可以认为在保持所有其他值不变的情况下，随着相关系数的降低，分散化所带来的潜在好处将会越来越大。

⊖ 如果相关系数为 +1，那么 $R_1 = a + bR_2$ ，其中 $b > 0$ 。

⊖ 对于正的相关系数（在 0 和 1 之间），当相关系数小于风险较小的那个资产的标准差除以风险较大的那个标准差时，就会给出一段斜率为负的部分。在我们的例子中，该比率等于长期政府债券的标准差和大盘股的标准差之比，或者 $10/15 = 0.6667$ 。因为 0.5 小于 0.6667，所以在相关系数为 0.5 情况下的最小方差前沿具有一段斜率为负的部分。这里，我们不允许卖空（负的资产权重）。如果我们允许卖空，那么对于正相关系数情况下的最小方差前沿都将具有斜率为负的部分，这可能涉及卖空风险更大的资产。更多相关的内容，参见 Elton, Gruber, Brown 和 Goetzmann (2003)。

- 当相关系数是 -1 时，最小方差前沿将由 2 条直线部分组成。这 2 个部分的交点为全局最小方差组合，该组合的方差为 0。在相关系数为 -1 的情况下，投资组合的风险可以被降到 0（如果我们想这么做的话）。
- 在 2 个极端相关系数 +1 和 -1 之间的最小方差前沿具有一个类似于子弹的形状。因此，最小方差前沿有时也被称为“子弹”（bullet）。
- 有效前沿是最小方差前沿中斜率为正的部分。在保持其他值不变的情况下，随着我们降低相关系数，有效前沿会得到改善，即在给定可以达到的收益率标准差的情况下，期望收益率能得到提高。

总之，当 2 个投资组合之间的相关系数小于 +1 时，分散化会带来潜在的好处。而在保持其他值不变的条件下，随着我们向 -1 的方向降低相关系数，分散化的潜在好处会增加。

11.2.2 扩展到 3 种资产的情况

之前我们已经讨论了如何形成包含 2 种资产（大盘股和政府债券）的投资组合。对于我们例子中想要在给定风险水平下最大化期望收益率的投资者（持有一个有效投资组合）来说，除非投资组合完全配置在股票之上，否则包含 2 种资产最优投资组合中的每种资产都会包含有一定的比例。

现在我们可能会问，如果将另一种资产加入到这个可能的投资选择中是否会改善我们这个可以获得的投资组合在风险和收益上的权衡？对于该问题的回答基本上是肯定的。一个基本的经济学原理告诉我们，一个额外的选择不会使我们的状况变差。在最差的情况下，一个投资者可以忽视额外的选择，而不会产生比之前更差的结果。然而，通常一个新的资产会使得我们移动到更好的最小方差前沿上。我们可以通过将 2 种资产（这里是大盘股和政府债券）的最小方差前沿和 3 种资产（大盘股，政府债券和小盘股）的最小方差前沿进行对比来给出这个常见的结果。

在表 11-1 所给出的我们最初 2 个资产的例子中，我们假定了大盘股和政府债券的期望收益率，方差和相关系数。现在假设我们有一个额外的投资选择（小盘股）。那么我们是否能够达到比只在 2 种资产（大盘股和政府债券）中选择更优的风险收益的权衡？

表 11-5 给出了我们对于这 3 种资产期望收益率以及收益率的标准差和它们之间相关系数的假设。

表 11-5 假定的期望收益率、方差和相关系数：3 种资产的例子			
	资产 1 大盘股	资产 2 政府债券	资产 3 小盘股
期望收益率	15%	5%	15%
方差	225	100	225
标准差	15%	10%	15%
相关系数			
大盘股和债券		0.5	
大盘股和小盘股		0.8	
债券和小盘股		0.5	

现在我们考虑这些统计量和投资组合期望收益率和方差之间的关系。对于任一个由 3 种资产组成的投资组合，假设投资组合中 3 种资产的权重分别为 $\omega_1, \omega_2, \omega_3$ ，投资组合的期望收益率 $E(R_p)$ 为

$$E(R_p) = \omega_1 E(R_1) + \omega_2 E(R_2) + \omega_3 E(R_3)$$

式中， $E(R_1)$ 表示资产 1（这里是大盘股）的期望收益率

$E(R_2)$ 表示资产 2（这里是政府债券）的期望收益率

$E(R_3)$ 表示资产 3（这里是小盘股）的期望收益率

投资组合的方差为

$$\sigma_p^2 = \omega_1^2 \sigma_1^2 + \omega_2^2 \sigma_2^2 + \omega_3^2 \sigma_3^2 + 2\omega_1 \omega_2 \rho_{1,2} \sigma_1 \sigma_2 + 2\omega_1 \omega_3 \rho_{1,3} \sigma_1 \sigma_3 + 2\omega_2 \omega_3 \rho_{2,3} \sigma_2 \sigma_3$$

式中， σ_1 表示资产 1 收益率的标准差
 σ_2 表示资产 2 收益率的标准差
 σ_3 表示资产 3 收益率的标准差
 $\rho_{1,2}$ 表示资产 1 和资产 2 之间的相关系数
 $\rho_{1,3}$ 表示资产 1 和资产 3 之间的相关系数
 $\rho_{2,3}$ 表示资产 2 和资产 3 之间的相关系数
投资组合的标准差为

$$\sigma_p = [\omega_1^2 \sigma_1^2 + \omega_2^2 \sigma_2^2 + \omega_3^2 \sigma_3^2 + 2\omega_1 \omega_2 \rho_{1,2} \sigma_1 \sigma_2 + 2\omega_1 \omega_3 \rho_{1,3} \sigma_1 \sigma_3 + 2\omega_2 \omega_3 \rho_{2,3} \sigma_2 \sigma_3]^{1/2}$$

给定我们的假设，投资组合的期望收益率为

$$E(R_p) = \omega_1(0.15) + \omega_2(0.05) + \omega_3(0.15)$$

投资组合的方差为

$$\begin{aligned} \sigma_p^2 = & \omega_1^2 0.15^2 + \omega_2^2 0.10^2 + \omega_3^2 0.15^2 + 2\omega_1 \omega_2 (0.5)(0.15)(0.10) \\ & + 2\omega_1 \omega_3 (0.8)(0.15)(0.15) + 2\omega_2 \omega_3 (0.8)(0.15)(0.15) \end{aligned}$$

投资组合的标准差为

$$\begin{aligned} \sigma_p = & [\omega_1^2 0.15^2 + \omega_2^2 0.10^2 + \omega_3^2 0.15^2 + 2\omega_1 \omega_2 (0.5)(0.15)(0.10) \\ & + 2\omega_1 \omega_3 (0.8)(0.15)(0.15) + 2\omega_2 \omega_3 (0.8)(0.15)(0.15)]^{1/2} \end{aligned}$$

然而，在这个 3 种资产的例子中，确定出最优的投资组合比 2 种资产的情况更难。在 2 种资产的例子中，大盘股在投资组合中所占的比例是 100% 减去政府债券所占的比例。但是在 3 种资产的例子中，我们需要一种方法来确定哪一个投资组合将在任一个给定期望收益率下得到最小的方差。我们首先至少知道最小期望收益率（将所有资产投入政府债券所得到的收益率 5%）以及最大期望收益率（没有资产投入政府债券所得到的收益率 15%）。而对于最大和最小期望收益率间的任一个期望收益率水平，我们必须求解出在该期望水平下获得最小风险的投资组合的权重。我们可以使用最优化程序（optimizer）（一种专门的计算机程序或者具有该功能的电子表格）来给出这些权重。[⊖]

注意到新资产（小盘股）具有与大盘股和债券都小于 +1 的相关系数，这表明小盘股可能在分散风险上非常有用。

正如表 11-6 所示，投资组合中大盘股和小盘股的比例在所有最小方差投资组合中都是相同的。这个比例源于表 11-5 中的简化假设，即大盘股和小盘股的收益率具有相同的期望值和标准差，以及相同的与政府债券收益率之间的相关系数。在一个不同的也更加现实的期望收益率、方差和相关系数的情况下，该例子的最小方差投资组合将包含不同的大盘股和小盘股的比例，但是我们仍将得到类似的获得改进可获得的风险—收益权衡的可能性。

表 11-6 对于 3 种资产情况下最小方差前沿上的点

期望 收益率(%)	投资组合 的方差	投资组合的 标准差(%)	大盘股 (ω_1)	政府债券 (ω_2)	小盘股 (ω_3)	期望 收益率(%)	投资组合 的方差	投资组合的 标准差(%)	大盘股 (ω_1)	政府债券 (ω_2)	小盘股 (ω_3)
5	100.00	10.00	0	1.00	0	11	124.90	11.18	0.30	0.40	0.30
6	96.53	9.82	0.05	0.90	0.05	12	139.73	11.82	0.35	0.30	0.35
7	96.10	9.80	0.10	0.80	0.10	13	157.60	12.55	0.40	0.20	0.40
8	98.72	9.94	0.15	0.70	0.15	14	178.53	13.36	0.45	0.10	0.45
9	104.40	10.22	0.20	0.60	0.20	15	202.50	14.23	0.50	0	0.50
10	113.13	10.64	0.25	0.50	0.25						

⊖ 这些程序使用被称为二次规划的求解方法。

3 种资产情形下每个期望收益率水平下的最小方差是如何与 2 种资产情形下的结果进行比较的呢? 图 11-5 给出了它们的一个比较。

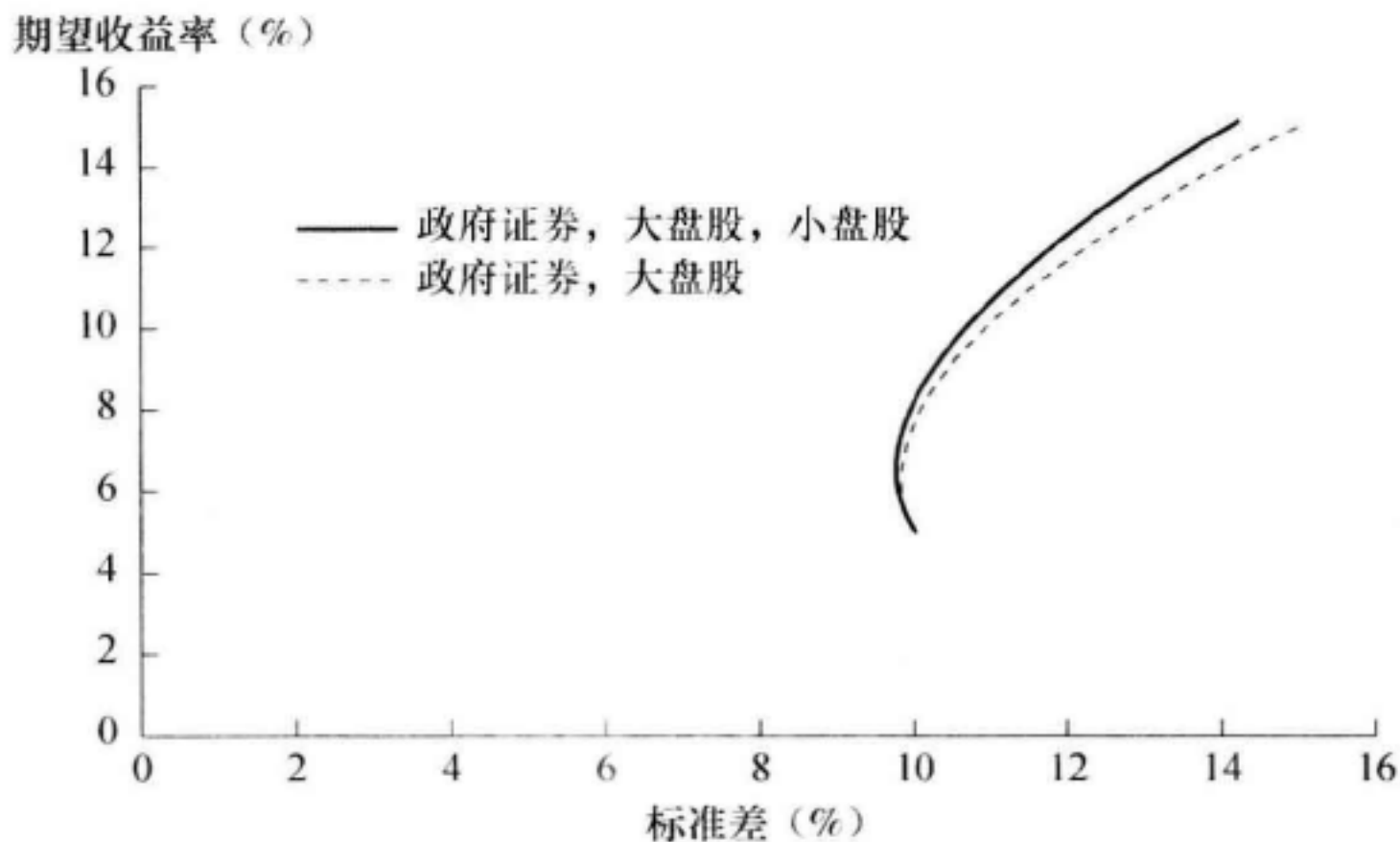


图 11-5 3 种资产与 2 种资产的最小方差前沿比较

当 100% 的投资组合投资到政府债券时, 2 种情况下的最小方差投资组合具有相同的期望收益率 (5%) 和标准差 (10%)。然而, 对于每个其他的期望收益率水平, 3 种资产情形下的最小方差投资组合比 2 种资产情况下相同期望收益率水平下的最小方差投资组合具有更小的标准差。同样我们也注意到 3 种资产的有效前沿优于 2 种资产的有效前沿 (我们将从最优的有效前沿上来选择最优的投资组合)。

从这个 3 种资产的例子中, 我们可以得到 2 个关于投资组合分散化理论的结论。首先, 我们一般可以通过扩大可以投资的资产集合来改善风险—收益权衡的结果。其次, 对于任一期望收益率水平下最小方差投资组合的构成依赖于期望收益率、收益率的方差、收益率之间的相关系数以及资产的数目。

11.2.3 多个资产最小方差前沿的确定

我们已经给出了 2 种和 3 种资产情况下均值—方差分析的例子。然而, 通常投资组合经理会使用大量的资产来建立最优投资组合。在这一节中, 我们将介绍如何来确定由许多资产所组成的投资组合的最小方差前沿。

对于由 n 种资产组成的投资组合, 投资组合的期望收益率为[⊖]

$$E(R_p) = \sum_{j=1}^n \omega_j E(R_j) \quad (11-1)$$

投资组合收益率的方差为[⊗]

$$\sigma_p^2 = \sum_{i=1}^n \sum_{j=1}^n \omega_i \omega_j \text{Cov}(R_i, R_j) \quad (11-2)$$

在确定最优投资组合权重之前, 记住投资组合中所有单个资产权重之和必须等于 1:

$$\sum_{j=1}^n \omega_j = 1$$

⊖ 求和符号表示我们将 j 从 1 变化到 n , 并加总这些由给定的 j 所得到的项。

⊗ 双重求和符号表明我们先将 i 设定为 1, 然后将 j 从 1 变化到 n , 再在将 i 设定为 2, 将 j 从 1 变化到 n , 以此类推, 直到 i 等于 n 未知; 然后将所有这些项进行加总。

为了确定 n 种资产集合所形成的最小方差前沿, 我们首先要确定该资产集合中可能的最小和最大的期望收益率 (这些收益率是指单个资产最小的期望收益率 $r_{\min}$ 和最大的期望收益率 $r_{\max}$)。然后我们必须确定在 $r_{\min}$ 和 $r_{\max}$ 之间每个期望收益率水平上所产生的最小方差投资组合中每个资产的权重。用数学术语来说, 我们必须在给定某个特定 z 值 ($r_{\min} \leq z \leq r_{\max}$) 的情况下, 求解如下问题:

$$\underset{\text{根据}\omega\text{的选择}}{\text{最小化}} \sigma_p^2 = \sum_{i=1}^n \sum_{j=1}^n \omega_i \omega_j \text{Cov}(R_i, R_j)$$

(11-3)

约束于 $E(R_p) = \sum_{j=1}^n \omega_j E(R_j) = z$ 和 $\sum_{j=1}^n \omega_j = 1$ 。

该最优化问题表明我们要求解投资组合的权重 ($\omega_1, \omega_2, \omega_3, \dots, \omega_n$) 以使得在给定期望收益率 z 水平和权重加总为 1 的约束条件下收益率的方差达到最小。该权重定义了一个投资组合, 而该投资组合是在其期望收益率水平下的最小方差组合。式 (11-3) 给出了只有投资组合权重加总等于 1 这个约束条件的最简单情况, 所以这个例子允许卖空资产。而对于限制卖空的约束将需要增加一个额外的假设 $\omega_j \geq 0$ 。我们通过将期望收益率从最小水平变动到最大水平来找出最小方差前沿。例如, 我们可以通过从 $z = r_{\min}$ 开始, 然后逐步对 z 增加 10 个基点 (0.10%) 直到求解出 $z = r_{\max}$ 的最优投资组合的权重, 来确定出小集合 z 值的最优投资组合的权重。[⊖] 我们使用一个最优化程序来实际求解最优化问题。例 11-2 给出了使用非美国股票和 3 种美国资产类型的历史数据所计算出的最小方差前沿。

例 11-2 使用国际的历史收益率数据所求得的最小方差前沿

在这个例子中, 我们使用 4 种资产类别来考察一个历史的最小方差前沿。3 种美国的资产类型分别是标准普尔 500 指数、美国小盘股以及美国长期政府债券。我们在这些资产类的基础上增加非美国股票 (MSCI 世界 (除美国外) 指数)。我们基于 1970 年 1 月到 2002 年 12 月的历史收益率数据考察了最小方差前沿。表 11-7 给出了这 4 种资产在整个样本区间的平均收益率、方差和相关系数。

表 11-7 4 种资产类型的平均年收益率, 标准差和相关系数矩阵 (1970 年 1 月到 2002 年 12 月)

	标准普尔 500	美国小盘股	MSCI 世界 (除美国外) 指数	美国长期债券
平均年收益率	11.6%	14.6%	11.1%	19.6%
标准差	15.83%	22.18%	17.07%	10.44%
相关系数				
标准普尔 500	1			
美国小盘股	0.731	1		
MSCI 世界 (除美国外) 指数	0.573	0.475	1	
美国长期债券	0.266	0.138	0.155	1

资料来源: Ibbotson Associates.

表 11-7 表明这 4 种资产类型的最小平均历史收益率为每年 9.6% (债券) 和最大平均历史收益率为 14.6% (美国小盘股)。为了找出最小方差前沿, 我们使用了最优化模型。该带有卖空限制的最优化程序使用下列等式计算出了均值—方差前沿:

⊖ 在不存在卖空的情况下存在着一个简单的求解方法。根据 Black (1972) 的两基金定理, 如果允许卖空, 所有在关于风险资产的最小方差前沿上的投资组合都是由任意 2 个最小方差投资组合的线性组合而形成的。这意味着如果我们已经计算出了 2 个最小方差投资组合中的权重, 那么我们就可以找出最小方差前沿。然而, 即使当我们增加许多投资者所面对的不能卖空的约束, 本文中所用的方法也是成立的。

$$\begin{aligned} \text{Min}\sigma_p^2(R) = & \omega_1^2\sigma_1^2 + \omega_2^2\sigma_2^2 + \omega_3^2\sigma_3^2 + \omega_4^2\sigma_4^2 + 2\omega_1\omega_2\rho_{1,2}\sigma_1\sigma_2 + 2\omega_1\omega_3\rho_{1,3}\sigma_1\sigma_3 + 2\omega_1\omega_4\rho_{1,4}\sigma_1\sigma_4 \\ & + 2\omega_2\omega_4\rho_{2,3}\sigma_2\sigma_3 + 2\omega_2\omega_4\rho_{2,4}\sigma_2\sigma_4 + 2\omega_3\omega_4\rho_{3,4}\sigma_3\sigma_4 \end{aligned}$$

受约束于 $E(R_p) = \omega_1 E(R_1) + \omega_2 E(R_2) + \omega_3 E(R_3) + \omega_4 E(R_4) = z$ (重复特定的 z 值, $0.096 \leq z \leq 0.146$), $\omega_1 + \omega_2 + \omega_3 + \omega_4 = 1$ 以及 $\omega_j \geq 0$ 。

权重 $\omega_1, \omega_2, \omega_3, \omega_4$ 分别代表在表 11-7 依次中列出的 4 种资产类型。随着我们从最小水平 ($r_{\min} = 9.6\%$) 变动到最大水平 ($r_{\max} = 14.6\%$), 最优化程序会选择产生每个给定平均收益率水平下的最小方差投资组合中的权重 (分配到 4 种资产类型)。在这个例子中, $E(R_j)$ 代表资产类 j 的样本平均收益率, 而方差和协方差同样也是样本统计量。除非我们故意选择使用这些历史数据作为我们对于未来的估计值, 否则我们会将这些最优化的结果解释为对于未来的预测值。

图 11-6 给出了基于 1970 ~ 2002 年历史的均值, 方差和协方差的 4 种资产类型的最小方差前沿。该图也分别给出了这 4 种资产类型的均值和标准差。

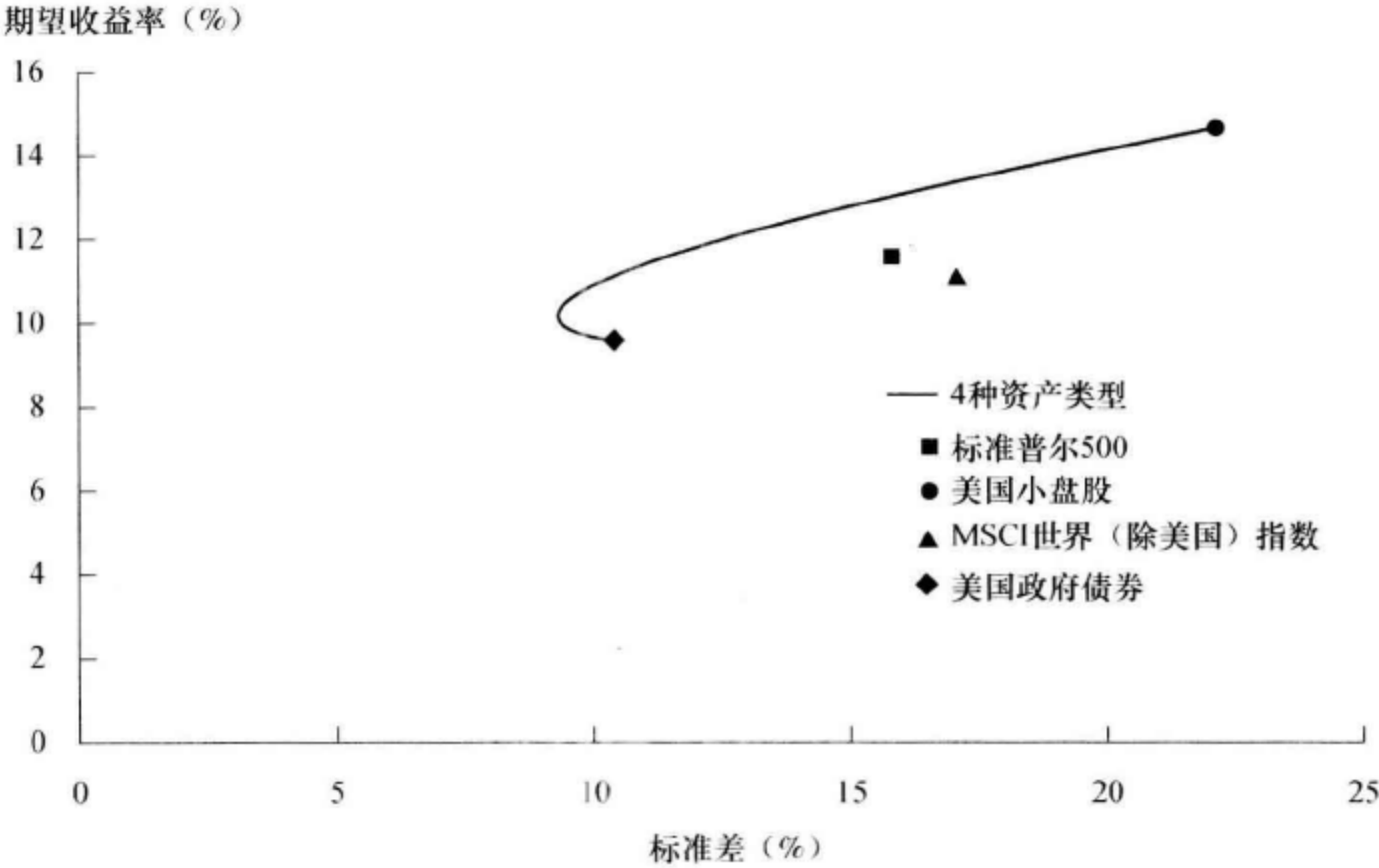


图 11-6 4 种资产类型的最小方差前沿 (1970 ~ 2002 年)

虽然美国政府债券位于最小方差前沿上, 但是它们被其他在相同风险水平上提供更好平均收益率的资产类型所占优。注意代表标准普尔 500 和 MSCI 世界 (除美国) 指数的点位于最小方差前沿的右下方。如果我们从标准普尔 500 或者 MSCI 世界 (除美国) 股票投资组合向左侧移动, 那么我们会到达有效前沿上的投资组合, 该投资组合在不影响平均收益率的情况下具有更小的风险。至少在了解到这个事实之后, 这 2 个资产组合对于投资于所有 4 种资产类型的投资者就不是有效的。尽管事实上 MSCI 世界 (除美国) 股票指数本身就是一个广泛市场的指数, 但是进一步的分散化仍旧对我们有利。在这 4 种资产类型之中, 只有美国小盘股是作为平均收益率最高的投资组合被绘制在有效前沿上的; 一般, 平均收益率最高的投资组合是在一个带有卖空约束的最优化中作为有效前沿的端点出现的, 正如在这个例子中所表现出来的一样。

在本节中，我们给出了求解最小方差前沿的过程。我们也对于由实际数据产生的前沿进行了分析。在下一节中，我们将讨论投资组合规模和分散化之间的关系。

11.2.4 分散化和投资组合的规模

之前，我们说明了将第3种资产增加到包含2种资产的投资组合中会带来分散化的好处。那个讨论开启了我们在本节中将要讨论的一个在实务中令人感兴趣的问题。到底我们需要持有多少只不同的股票才能产生一个充分分散化的投资组合？在确定一个投资组合的风险时，协方差或者相关系数是如何与投资组合的规模相互影响的？

我们使用一个持有相同资产权重投资组合的投资者的例子来回答这个问题。假如我们购买一个包含 n 种股票的投资组合，并在投资组合的每种股票中投资相同的比例 ($\omega_i = 1/n$, $i = 1, 2, \dots, n$)。收益率的方差为

$$\sigma_p^2 = \sum_{i=1}^n \sum_{j=1}^n \omega_i \omega_j \text{Cov}(R_i, R_j) = \frac{1}{n^2} \sigma_p^2 \sum_{i=1}^n \sum_{j=1}^n \text{Cov}(R_i, R_j) \quad (11-4)$$

假如我们将所有股票收益率的平均方差记为 $\bar{\sigma}^2$ ，而将所有两两配对股票收益率的平均协方差记为 $\overline{\text{Cov}}$ 。那么式 (11-4) 就可能简化为如下的形式^①：

$$\sigma_p^2 = \frac{1}{n} \bar{\sigma}^2 + \frac{n-1}{n} \overline{\text{Cov}} \quad (11-5)$$

随着股票数量 n 的增加，单只股票收益率方差的贡献会变得很小，因为 $\frac{1}{n} \bar{\sigma}^2$ 随着 n 的增大，其极限为 0。同样地，所有股票平均协方差对于投资组合方差的贡献将保持为一个非零的数，因为 $\frac{n-1}{n} \overline{\text{Cov}}$ 随着 n 的增加将收敛到 $\overline{\text{Cov}}$ 。因此，随着投资组合中资产数目的增加，投资组合的方差大约会等于平均协方差。在较大的投资组合中，平均协方差（刻画了资产是如何一起变动的），成了比单个风险或方差的平均值更加重要的因素。

除了这点之外，式 (11-5) 还使得我们度量出了从只持有单只股票的完全不分散的头寸到持有分散化投资组合的收益率方差的减少程度。如果一个投资组合只包含单只股票，那么它的方差当然就会是单只股票的方差，即持有最大方差的头寸。^② 而如果该投资组合包含非常多的股票，那么该投资组合的方差就将接近于两两股票间的平均协方差，其被称为持有最小方差的头寸。这两个方差水平到底差别有多大呢？我们在用相对较少数量股票进行的分散化投资中可以获得的最大好处是多少？

这些问题的回答依赖于平均方差和平均协方差的大小。因为相关系数要比协方差容易解释，所以我们将对相关系数进行考察。为了简单起见，假如任意 2 只股票的收益率之间的相关系数是相同的，并且所有股票具有相同的标准差。陈 (Chan)、卡塞斯基 (Karceski) 和兰考尼肖科 (Lakonishok) (1999) 发现对于 1968 ~ 1998 年时段的在美国 NYSE 和 Amex 中股票，小盘股收益率的相关系数为 0.24，大盘股收益率的相关系数为 0.33，而整个股票样本中股票收益率的平均相关系数为 0.28。假设共同的相关系数为 0.30，该值位于许多时间段美国股票的平均相关系数的估计区间之中。2 个随机变量的协方差是这 2 个变量的相关系数乘以它们的标准差，所以 $= 0.30 \sigma^2$ （使用我们的假设：所有股票的收益率都具有相同的标准差，

① 参见 Bodie, Kane 和 Marcus (2001)。

② 对于相关系数取实际值的情况，平均方差要大于平均协方差。

记为 σ)。

再看一下式 (11-5), 并将 $\overline{\text{Cov}}$ 用 $0.30\sigma^2$ 代替:

$$\begin{aligned}\sigma_p^2 &= \frac{1}{n}\sigma^2 + \frac{n-1}{n}(0.30\sigma^2) = \frac{\sigma^2}{n}[1 + 0.30(n-1)] \\ &= \frac{\sigma^2}{n}(0.70 + 0.30n) = \sigma^2\left(\frac{0.70}{n} + 0.30\right)\end{aligned}$$

该式给出了一个更加一般的表达式 (假设股票的收益率具有相同的标准差):

$$\sigma_p^2 = \sigma^2\left(\frac{1-\rho}{n} + \rho\right) \quad (11-6)$$

如果投资组合只包含一只股票, 那么投资组合的方差就是 σ^2 。随着 n 的增加, 投资组合的方差会迅速下降。在我们例子中, 如果投资组合包含 15 只股票, 那么投资组合的方差将为 $0.347\sigma^2$, 其只是包含单只股票投资组合的方差的 34.7%。在 30 只股票的情况下, 投资组合的方差是单只股票投资组合方差的 32.3%。这个例子中最小的可能投资组合的方差为单只股票方差的 30%, 因为当 n 趋近于无穷大时, $\sigma_p^2 = 0.30\sigma^2$ 。而只要具有 30 只股票, 投资组合的方差就只比最小可能值大 8% 左右 ($0.323\sigma^2/0.30\sigma^2 - 1 = 0.077$), 而比包含单只股票的投资组合的方差要小 67.7%。

通过对于一个相关系数合理的假定值, 前面的例子可以表明一个由许多股票组成的投资组合具有比由单只股票组成的投资组合更小的总风险。在这个例子中, 我们可以通过持有许多股票来分散 70% 的单只股票的风险。而且, 我们可以通过持有非常少的证券来获得大部分分散化所带来的风险降低的好处。

那么, 如果股票间的相关系数高于 0.30, 又会怎样呢? 假如一个投资者想要使其投资组合的风险仅是最小可能分散化的投资组合风险的 110%, 那么该投资者需要持有多少股票? 如果股票间平均的相关系数为 0.5, 那么他将只需要用 10 只股票来形成投资组合以使其风险为最小可能投资组合方差的 110%。当相关系数更高时, 投资者将需要持有更少的股票来达到该风险水平。那么如果相关系数低于 0.30 时, 又会怎么样呢? 如果股票间的相关系数为 0.1, 那么投资者需要持有 90 只股票来达到该风险水平。

投资者的一个共同认识是几乎所有的分散化好处都可以由一个只包含 30 只股票的投资组合来实现。事实上, 费雪 (Fisher) 和劳丽 (Lorie) (1970) 表明 1926 ~ 1965 年纽交所交易股票的 95% 的分散化好处可以通过一个由 32 只股票组成的投资组合来实现。

正如上面所示, 实现某个水平分散化好处所需的股票数量依赖于股票收益率间的平均相关系数: 平均相关系数越低, 所需的股票数量越多。坎贝尔 (Campbell)、列图 (Lettau), 麦基尔 (Malkiel) 和徐 (Xu) (2001) 表明虽然总体市场波动自 1963 年后就没有再增长, 但是个股收益率在最近 (1986 ~ 1997 年) 的波动更加剧烈, 而个股收益率间的相关性也更弱了。因此, 在距离现在更近的时间段要实现相同风险降低的好处需要在投资组合中包含比费雪和劳丽所研究时间段更多的股票。坎贝尔等人认为在 1963 ~ 1985 年间, “一个包含 20 只股票的投资组合会降低年化超额标准差约 5%, 但是在 1986 ~ 1997 年间的子样本中, 降低到该超额标准差需要近 50 只股票。”^①

① Campbell 等人将“超额标准差”定义为给定大小的随机选取投资组合的标准差减去等权重市场指数的标准差。

例 11-3 伯克郡哈撒韦公司的分散化

伯克郡哈撒韦公司非常成功的首席执行官沃伦·巴菲特是一个对现代投资组合理论和分散化提出最严厉批评的反对者之一。例如，巴菲特曾经说过：“如果你是一位知情的投资者，而且能够了解商业经济，并找出 5~10 个定价合理且拥有重要的长期竞争优势的公司，那么传统的分散化将对你毫无用处，因为它将会损害你的收益结果并增加你的风险。”[Ⓐ]

那么巴菲特是否完全避免分散化呢？从他的投资记录上来看确实是这样的，但是即使是巴菲特也在一定程度上从事于分散化的投资。例如，考察伯克郡哈撒韦公司在 2002 年年末持有的三大投资。[Ⓑ]

美国运通公司	\$56 亿 (32%)
可口可乐公司	\$88 亿 (51%)
吉利公司	\$29 亿 (17%)
总投资	\$173 亿

这 3 只股票提供了多大的风险分散化？该投资组合的标准差比只包含可口可乐公司股票投资组合的标准差低多少？为了回答这些问题，假设这些股票的历史平均收益率、收益率的标准差和收益率之间的相关系数是未来的预期收益率、收益率标准差和相关系数的最佳估计。表 11-8 基于 1990~2002 年的月度收益率数据给出了这些历史统计量。

表 11-8 历史收益率、方差和相关系数：伯克郡哈撒韦公司持有量最大的股票（月度数据，1990 年 1 月~2002 年 12 月）

	美国运通	可口可乐	吉利
平均年收益率	16.0%	16.1%	17.6%
标准差	29.0%	24.7%	27.3%
相关系数			
美国运通和可口可乐		0.361	
美国运通和吉利		0.317	
可口可乐和吉利		0.548	

资料来源：FactSet.

表 11-8 表明在该时间段可口可乐公司股票的平均年收益率为 16.1%，年收益率的标准差为 24.7%。相反，包含美国运通公司股票（32%）、可口可乐公司股票（51%）和吉利公司股票（17%）的投资组合具有 $0.32(0.16) + 0.51(0.161) + 0.17(0.176) = 0.163$ 或者 16.3%。

基于表 11-8 中的权重和统计量，投资组合的期望标准差为

或者

$$\sigma_p = [\omega_1^2\sigma_1^2 + \omega_2^2\sigma_2^2 + \omega_3^2\sigma_3^2 + 2\omega_1\omega_2\rho_{1,2}\sigma_1\sigma_2 + 2\omega_1\omega_3\rho_{1,3}\sigma_1\sigma_3 + 2\omega_2\omega_3\rho_{2,3}\sigma_2\sigma_3]^{1/2}$$
$$\sigma_p = [(0.32^2)(0.290^2) + (0.51^2)(0.247^2) + (0.172)(0.273^2) + 2(0.32)(0.51)(0.361)(0.290)(0.247) + 2(0.32)(0.17)(0.317)(0.290)(0.273) + 2(0.51)(0.17)(0.548)(0.247)(0.273)]^{1/2} = 0.210 \text{ 或者 } 21.0\%$$

包含这 3 只股票投资组合的标准差是只包含可口可乐公司股票投资组合标准差的 $21.0/24.7 = 85.0\%$ 。因此，伯克郡哈撒韦公司实际上在风险降低的意义下实现了极大的分散化，甚至在只考虑持有最多的三大股票的情况下也是这样。

Ⓐ 参见 Buffett (1993)。
Ⓑ 我们只考察持有的三大投资是为了使得这个例子中的计算变得简单。同样为了简化，我们对投资组合中资产配置权重的百分比进行了四舍五入。这里给出的权重是在这 3 只股票投资中的相对权重，不是它们在伯克郡哈撒韦公司投资组合中的实际权重。

11.2.5 存在无风险资产条件下的投资组合选择

至今，我们只考虑了包含风险证券的投资组合，这隐含地表示投资者无法投资无风险资产。但是投资者可以持有政府证券，如国库券，而它们是在适当的时间段中在名义上近似无风险的。例如，购买1年期国库券的投资者了解其持有到期的名义收益率。那么在我们可以投资于无风险资产时的风险和收益率之间的权衡又是怎样的呢？

无风险资产收益率的标准差是0，因为收益率是确定的，并且不存在违约风险。假如无风险资产的收益率是每年4%。如果我们将国库券作为无风险资产，那么4%是事先已知的实际收益率；它不是一个预期收益率。[⊖]因为无风险资产收益率的标准差为0，无风险资产收益率和其他资产收益率之间的协方差也一定为0。这些观测能帮助我们理解增加一个无风险资产是如何影响资产间均值方差权衡的。

11.2.5.1 资本配置线

资本配置线（Capital allocation line, CAL）描述了投资者最优投资组合中风险资产和无风险资产可以获得的期望收益率和收益率标准差组合的全体。因此，CAL描述了投资者关于如何在风险资产和无风险资产间进行最优配置决策的期望结果。

在平均收益率—标准差空间中的图像是否满足CAL的定义？CAL一定是从无风险收益率出发的一条直线，该直线与风险资产的有效前沿相切；由无风险收益率出发到风险资产有效前沿上的所有直线中，CAL这条相切的直线具有最大的斜率和最好的风险—收益的权衡（“相切”是指接触而不相交）。切点投资组合是投资者最优的风险资产组合。投资者的风险容忍程度决定了他在该直线上所选择的点。例11-4以及紧接着的讨论阐述了这些内容。

例 11-4 存在无风险资产情况下投资者风险和收益之间的权衡

假如我们想要确定将无风险资产加入到原来大盘股和政府债券投资组合后所产生的效果。

表11-9给出了这3种资产的假定期望收益率、方差和相关系数。

假如我们决定将所有资产投资到收益率为4%的无风险资产之中。在这个例子中，该投资组合的期望收益率是4%，而期望标准差是0。现在假设我们将所有资产投资到大盘股中。那么期望收益率为15%，而投资组合的标准差为15%。如果我们将投资组合在无风险资产和大盘股之间进行分配，那么又会发生什么呢？如果投资于大盘股资产的比例为 ω_1 ，而投资于无风险资产的比例为 $(1 - \omega_1)$ ，那么期望投资组合收益率为

$$E(R_p) = \omega_1 (0.15) + (1 - \omega_1) (0.04)$$

而投资组合标准差为

表 11-9 期望收益率，方差和相关系数：包括无风险资产的3种资产例子

	大盘股	政府债券	无风险资产
期望收益率	15%	5%	4%
方差	225	100	0
标准差	15%	10%	0%
相关系数			
大盘股和政府债券		0.5	
大盘股和无风险资产		0	
政府债券和无风险资产		0	

⊖ 我们这里假设国库券的到期日与投资期限一样，因此，我们没有利率风险。

$$\sigma_p = [\omega_1^2(1.15)^2 + (1 - \omega_1)^2(0)^2]^{1/2} = \omega_1(0.15)$$

注意到期望收益率和收益率的标准差都是关于投资组合中投资于大盘股的比例 ω_1 呈线性的。图 11-7 给出了该例子对于无风险资产和大盘股之间风险和收益的权衡。

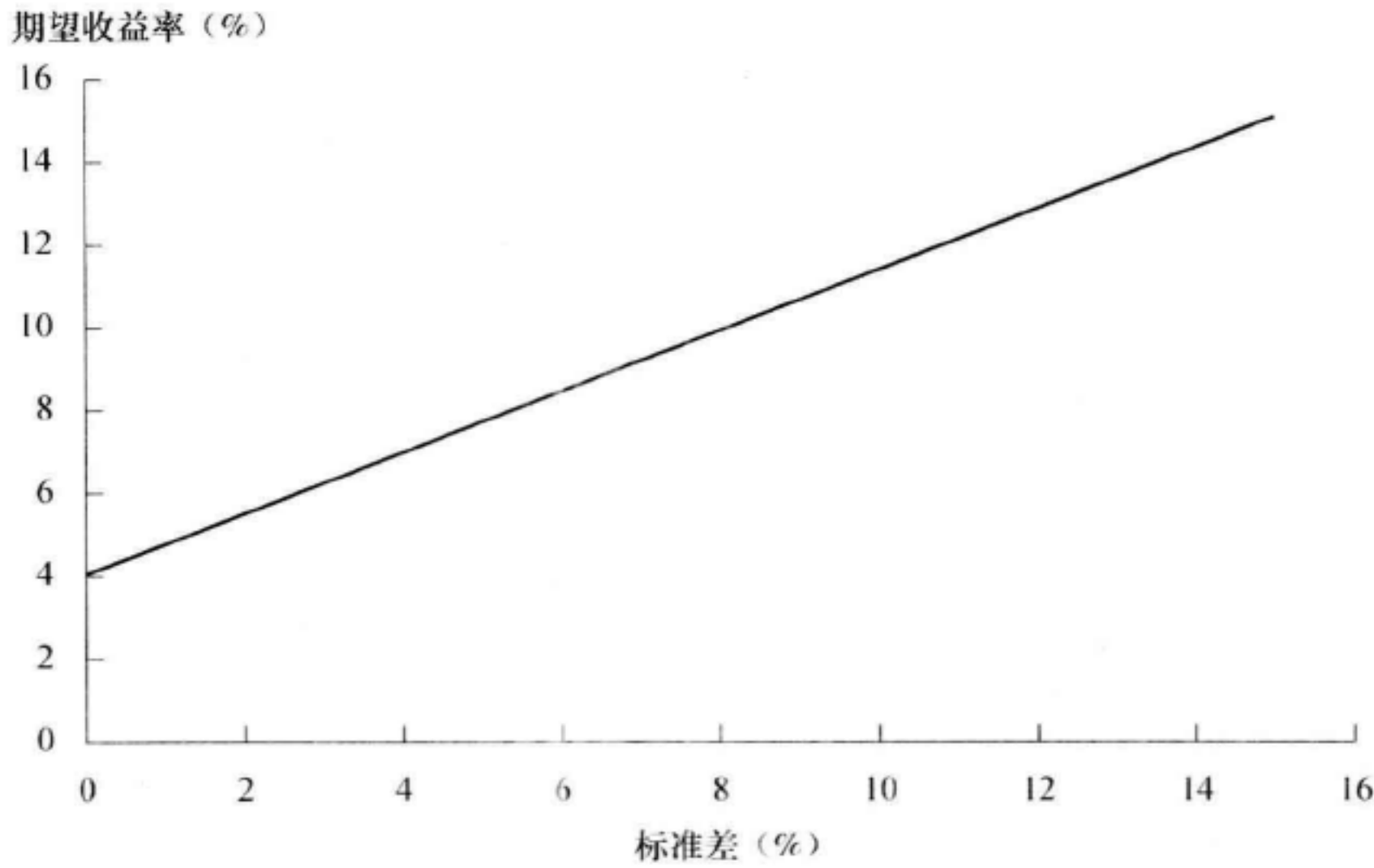


图 11-7 包含无风险资产和大盘股的投资组合

现在让我们考虑包含无风险资产和美国政府债券投资组合的风险和收益率之间的权衡。假如我们决定将所有资产投资到无风险资产中。在这个例子中，投资组合的期望收益率为 4%，而期望标准差为 0%。如果我们决定将所有资产投资到美国政府债券之中。那么投资组合的期望收益率为 5%，而投资组合的标准差为 10%。如果我们将投资组合在无风险资产和政府债券之间进行分配，那么又会发生什么呢？如果投资于政府债券的比例为 ω_1 ，而投资于无风险资产的比例为 $(1 - \omega_1)$ ，那么期望投资组合收益率为

$$E(R_p) = \omega_1(0.05) + (1 - \omega_1)(0.04)$$

而投资组合标准差为

$$\sigma_p = [\omega_1^2(1.10)^2 + (1 - \omega_1)^2(0)^2]^{1/2} = \omega_1(0.10)$$

注意到期望收益率和收益率的标准差都是关于投资组合中投资于政府债券的比例 ω_1 呈线性的。图 11-7 给出了该例子对于无风险资产和政府债券之间风险和收益的权衡。

我们刚才已经看到了 2 个不同的投资组合风险和收益之间的权衡：一个是无风险资产和大盘股的投资组合，而另一个是无风险资产和政府债券的投资组合。那么我们如何将这 2 个风险与收益的权衡与之前的政府债券和大盘股之间的风险—收益权衡进行比较呢？图 11-9 给出了所有这 3 个投资组合风险—收益的权衡。

注意到代表无风险资产和政府债券投资组合的直线与债券和股票的最小方差前沿交于债券—股票最小方差前沿最小收益率的一点，即投资组合中 100% 投资于债券的一点。代表无风险资产—债券直线上的点比债券—股票前沿上的点具有更小的风险和收益率；然而，我们不能在无风险资产—债券直线上找到在给定风险水平下，比债券—股票前沿更高期望收益率的点。在这个例子中，如果我们画出从无风险资产到债券—股票前沿的具有最高斜率的直线，那么该直

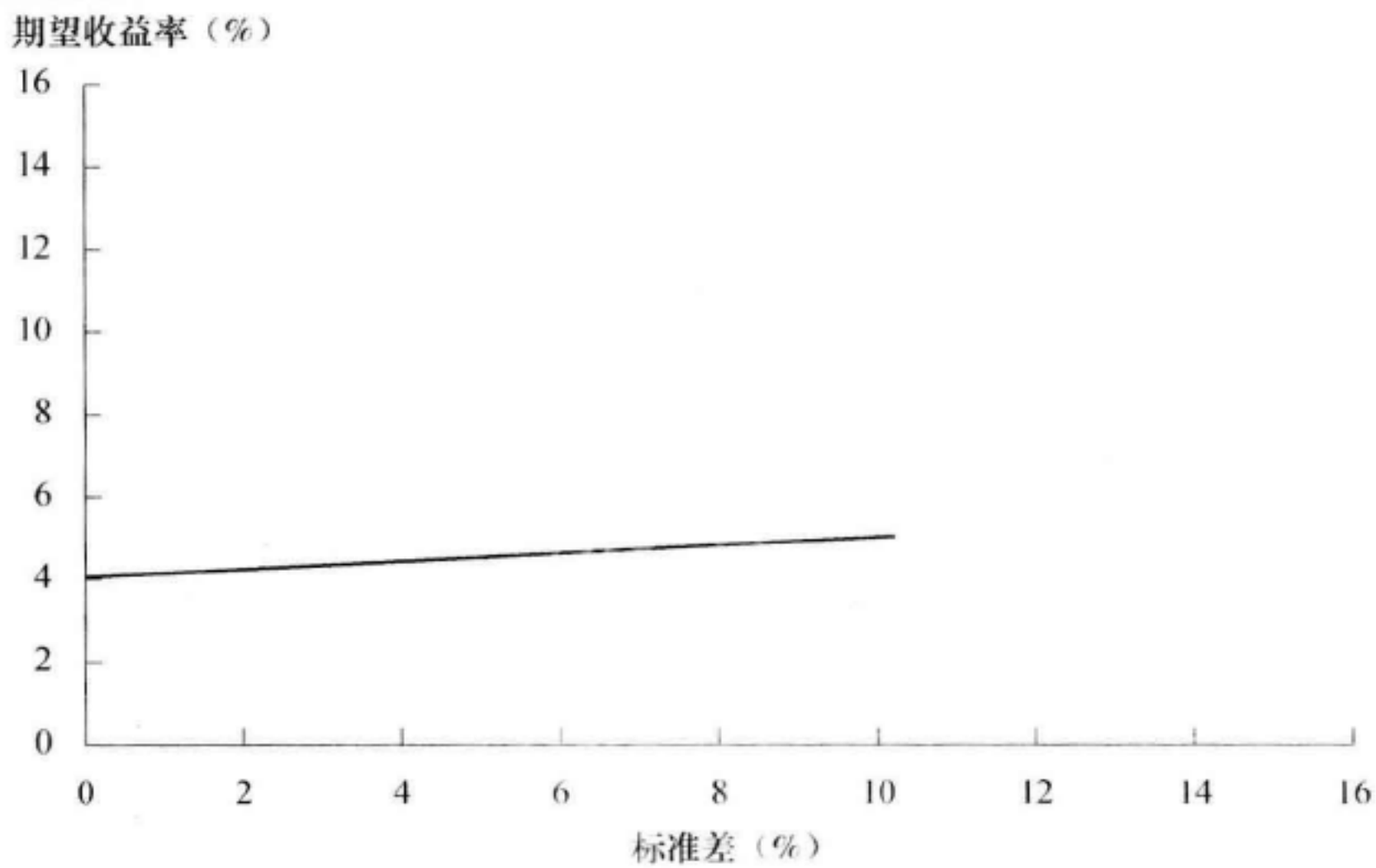


图 11-8 包含无风险资产和政府债券的投资组合

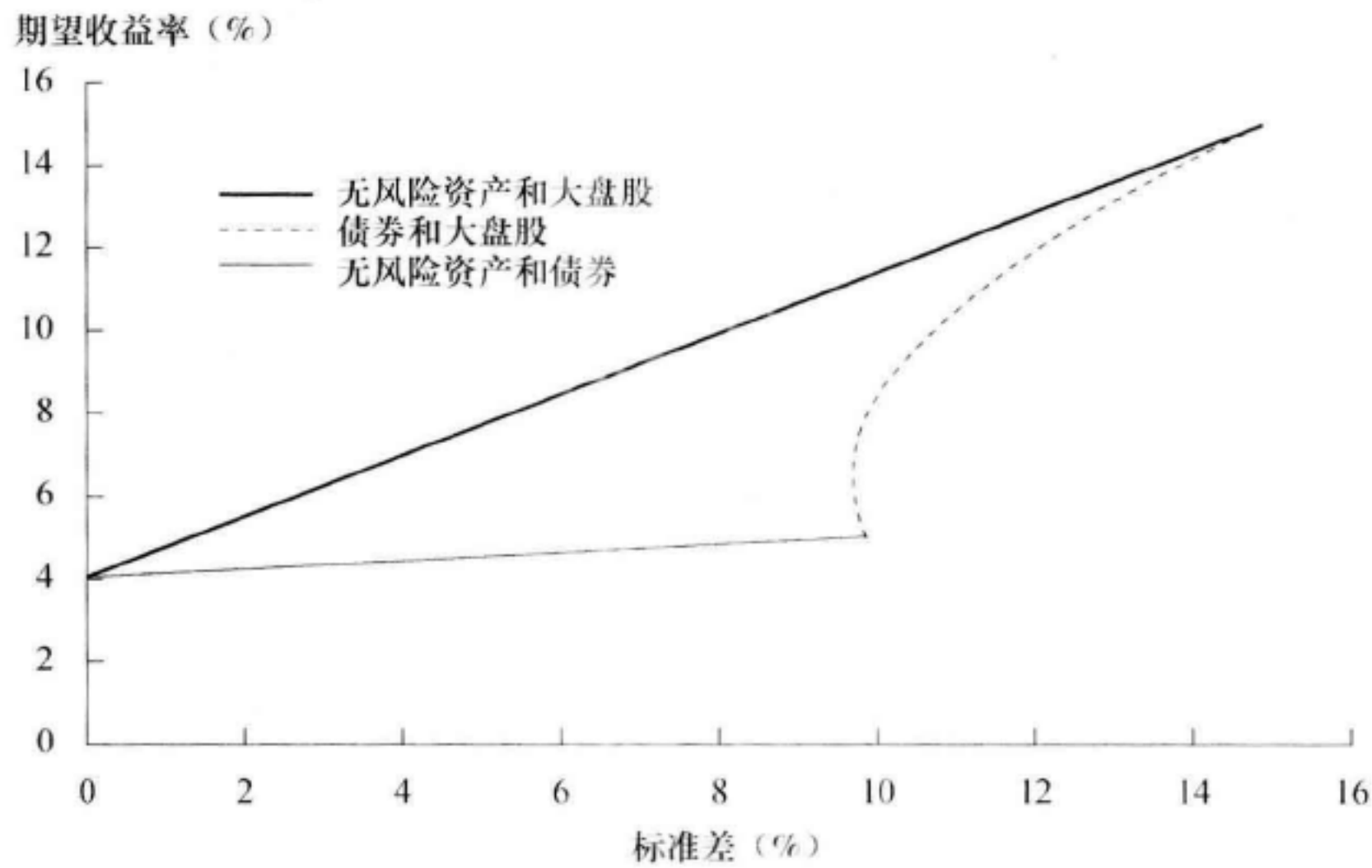


图 11-9 包含无风险资产，大盘股和政府债券的投资组合

线是与债券—股票前沿相切于代表 100% 投资于股票的一点。[⊖]该 CAL 以无风险资产和大盘股在图 11-9 中予以了标出。给定对于期望收益率、方差和协方差的假定，如果可以将我们的资金在最优风险投资组合和无风险资产间进行分配，那么资本配置线能够找出给定风险水平上具有最大期望收益率的投资组合。在所有从无风险收益率出发到风险资产的最小方差前沿的直线中，CAL 具有最大的斜率。斜率定义为上升的部分（期望收益率）与水平区间（标准差）的比率，它度量了期望风险和收益之间的权衡。CAL 是与最小方差前沿接触的斜率最大的直线；因此，给定我们的期望，资本配置线提供了能够获得的最佳的风险—收益权衡。

⊖ 然而，在包含许多资产的一般例子中，从无风险资产出发具有最大斜率的直线与风险资产最小方差前沿的交点并不代表只包含期望收益最高资产的投资组合。在这种情况下，CAL 代表无风险资产和许多风险资产的一个组合。

前面的例子给出了包含无风险资产的投资组合关于风险—收益权衡的 3 个重要的一般法则：

- 如果我们能够投资于无风险资产，那么 CAL 代表可以达到的最优的风险收益权衡。
- CAL 在 y 轴上的截距等于无风险收益率。
- CAL 与风险资产的有效前沿相切。[⊖]

例 11-5A 多种资产情况下的 CAL

在例 11-4 中，CAL 与所有风险资产形成的有效前沿相切。我们现在说明有效前沿是如何随着机会集（opportunity set）（可以进行投资的资产集合）和投资者是否想要借钱投资的情况的变化而发生变化的。我们可以通过再次考虑先前关于标准普尔 500、美国小盘股、非美国股票（MSCI 世界指数（不包括美国））、美国政府债券和无风险资产最优投资组合选择的例子（例 11-2）来说明这一点。

现在我们假设无风险利率为 5%。无风险收益率的标准差为 0，因为其收益率是确定的；无风险资产的收益率和其他资产收益率间的协方差也为 0。我们可以得到如下法则：

- 最大期望收益率的点不是 CAL 和风险资产有效前沿的切点。
- CAL 和风险资产有效前沿的切点代表只包含风险资产而没有无风险资产的投资组合。
- 如果我们排除以无风险利率借款的情况，那么所有资产（包括无风险资产）的有效前沿不可能是完全线性的。

图 11-10 给出了所有这 5 种资产的均值—方差前沿（最初的 4 种加上无风险资产）（假设不允许以无风险利率进行借款）。

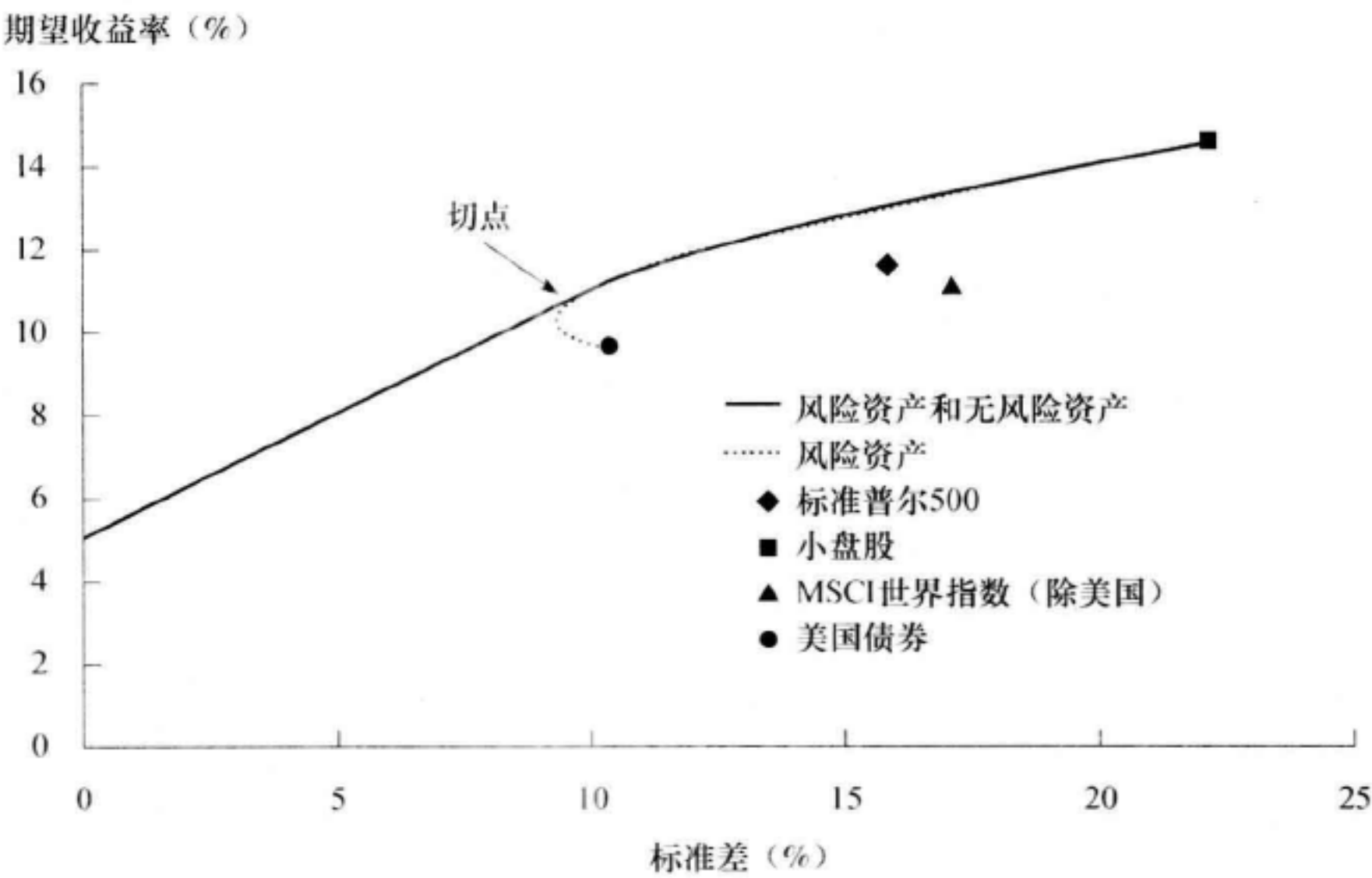


图 11-10 在风险资产投资组合中增加无风险资产后产生的影响

⊖ 注意到当我们将资产集扩展到包含无风险资产的资产集时，一个投资者的 CAL 就成为与扩展后的资产集相关的有效前沿。风险资产的有效前沿是只考虑风险资产的有效前沿。理解有效前沿的定义总是与特定的资产集有关是非常重要的。

正如该图所示,有效前沿从 y 轴截距(将所有资产投入无风险资产所形成的风险—收益组合)到切点是线性的。然而,如果投资者想要不通过借款在切点右侧获得更高的额外收益率,那么投资组合的有效前沿就将是切点右侧由4种风险资产所产生的有效前沿部分。投资者的有效前沿具有线性和曲线2个部分。^①

11.2.5.2 资本配置线方程

在前一节中,我们给出了CAL的图像,并讨论了在不同可能投资的资产集和投资组合经理预期条件下有效前沿的变化。我们现在给出CAL的方程。

假设一位投资者,给定对于风险资产均值、方差和协方差的预期,我们可以绘制出风险资产的有效前沿。存在一个提供无风险收益率 R_F 的无风险资产。如果 w_T 是投资者在投资组合中投资于切点组合的比例,那么整个投资组合的期望收益率为:

$$E(R_p) = (1 - w_T)R_F + w_TE(R_T)$$

而该投资组合的标准差为

$$\begin{aligned}\sigma_p &= [(1 - w_T)^2 \sigma_{R_F}^2 + w_T^2 \sigma_{R_T}^2 + 2(1 - w_T)w_T \sigma_{R_F} \sigma_{R_T} \rho_{R_F, R_T}]^{1/2} \\ &= [(1 - w_T)^2 (0) + w_T^2 \sigma_{R_T}^2 + 2(1 - w_T)w_T (0) \sigma_{R_T} (0)]^{1/2} \\ &= (w_T^2 \sigma_{R_T}^2)^{1/2} \\ &= w_T \sigma_{R_T}\end{aligned}$$

一个投资者可以选择投资到无风险资产与切点投资组合的比例;因此,他可以选择许多不同的风险与收益的组合。如果他全部投资到无风险资产之中,那么

$$\begin{aligned}w_T &= 0 \\ E(R_p) &= (1 - 0)R_F + 0E(R_T) = R_F \\ \sigma_p &= 0\sigma_{R_T} \\ &= 0\end{aligned}$$

如果他全部投资于切点组合,那么

$$\begin{aligned}w_T &= 1 \\ E(R_p) &= (1 - 1)R_F + 1E(R_T) = E(R_T) \\ \sigma_p &= 1\sigma_{R_T} = \sigma_{R_T}\end{aligned}$$

一般地,如果他投资于切点组合的比例为 w_T ,那么他投资组合的标准差将为 $\sigma_p = w_T \sigma_{R_T}$ 。

为了了解投资组合权重、期望收益率和风险之间的关系,我们使用关系式 $w_T = \sigma_p / \sigma_{R_T}$ 。如果我们将 w_T 的值回代到期望收益率的表达式 $E(R_p) = (1 - w_T)R_F + w_TE(R_T)$ 中,那么我们会得到

$$E(R_p) = \left(1 - \frac{\sigma_p}{\sigma_{R_T}}\right)R_F + \frac{\sigma_p}{\sigma_{R_T}}E(R_T)$$

或者

$$E(R_p) = R_F + \frac{E(R_T) - R_F}{\sigma_{R_T}} \sigma_p \quad (11-7)$$

该等式给出了给定投资者预期情况下可能的最佳预期风险与收益之间的权衡。 $\frac{E(R_T) - R_F}{\sigma_{R_T}}$ 这一

① 如果允许以无风险利率进行借贷(等价于在保证金账户以无风险收益率进行借款购买),有效前沿将是从小风险利率出发、穿过切点的一条直线。

项是投资者为了承担每一单位额外风险所要求的回报。例 11-5B 说明了如何利用资本配置线计算投资组合风险的价格以及与该投资者投资有关其他方面的问题。

例 11-5B CAL 的计算

假如无风险收益率 R_F 为 5%；投资者对于切点组合的预期收益率为 $E(R_T)$ ；并且切点组合的标准差为 25%。

(1) 该投资者对于每一单位额外的风险会要求多高的收益率？

(2) 假如该投资者需要一个收益率的标准差为 10% 的投资组合。那么他应该投资到切点投资组合的比例应该是多少，而此时该投资组合的预期收益率又是多少？

(3) 假如该投资者想要将投资组合中 40% 的资产投资到无风险资产之中。该投资组合的预期收益率将为多少？标准差又将为多少？

(4) 投资者对于标准差为 35% 的投资组合所要求的预期收益率为多少？

(5) 为了使得投资组合的预期收益率为 19%，投资者需要持有的切点组合和无风险资产的组合是怎样的？

(6) 如果该投资者具有 1 千万美元要投资，那么她必须以无风险收益率借入多少金额才能使得投资组合具有 19% 的预期收益率？

解：

(1) 在这个例子中， $[E(R_T) - R_F] / \sigma_{R_T} = (0.15 - 0.05) / 0.25 = 0.4$ 。投资者对于投资组合收益率标准差每一个百分比的增长，需要额外 40 个基点的期望收益率。

(2) 因为 $\sigma_p = w_T \sigma_{R_T}$ ，于是 $w_T = 0.1 / 0.25 = 0.4$ 或者 40%。换句话说，40% 的资产投资于切点组合，而 60% 投资于无风险资产。投资组合的期望收益率为 $E(R_p) = R_F + \frac{E(R_T) - R_F}{\sigma_{R_T}} \sigma_p = 0.05 + (0.4)(0.1) = 0.09$ 或者 9%。

(3) 在这个例子中， $w_T = 1 - 0.4 = 0.6$ 。因此，期望投资组合收益率为 $E(R_p) = (1 - w_T) R_F + w_T E(R_T) = (1 - 0.6)(0.05) + (0.6)(0.15) = 0.11$ 或者 11%。投资组合的标准差为 $\sigma_p = w_T \sigma_{R_T} = (0.6)(0.25) = 0.15$ 或者 15%。

(4) 我们知道该投资组合的风险和期望收益率之间的关系为 $E(R_p) = R_F + \frac{E(R_T) - R_F}{\sigma_{R_T}} \sigma_p = 0.05 + [(0.15 - 0.05) / 0.25] \sigma_p = 0.05 + 0.4 \sigma_p$ 。如果投资组合收益率的标准差为 35%，那么投资者所要求的期望收益率为 $E(R_p) = 0.05 + 0.4 \sigma_p = 0.05 + 0.4(0.35) = 0.19$ 或者 19%。

(5) 在期望收益率为 19% 的条件下，资产配置一定表示如下：

$$E(R_p) = (1 - w_T) R_F + w_T E(R_T) \text{ 或者}$$

$$0.19 = (1 - w_T)(0.05) + w_T(0.15) = 0.05 + 0.10w_T$$

$$w_T = 1.4$$

如何使得切点组合的权重为 140%？ $w_T = 1.4$ 可以解释为切点组合中的投资额，包括 (1) 全部初始财富的数量；(2) 以无风险利率借来的等于初始财富 40% 的金额。我们可以得到其期望收益率为 $(-0.4)(0.05) + (1.4)(0.15) = 0.19$ 或者 19%。

(6) 投资者必须以无风险利率借入 400 万美元来使切点资产组合的持有量增加到 1 400 万美元。因此投资组合的净值将为 $\$14\,000\,000 - \$4\,000\,000 = \$10\,000\,000$ 。

在这一节中,我们假设投资者可以对于风险资产的平均收益率、收益率的方差以及相关系数持有不同的看法。因此,每个投资者可能得到关于风险资产的不同有效前沿和不同的切点投资组合(投资者用来与无风险借贷资产进行组合的风险资产的最优投资组合)。在下面两节中,我们将考察在均值一方差投资者具有相同预期情况下的结果。

11.2.5.3 资本市场线

当投资者对于风险资产的平均收益率、收益率的方差以及相关系数具有相同预期时,所有投资者的CAL都是相同的,这时CAL被称为资本市场线(capital market line, CML)。在相同预期条件下,切点组合一定对于每个投资者来说都是相同的。在均衡时,切点组合一定是以市值加权的包含所有风险资产的一个投资组合;该切点组合是所有风险资产的市场组合。CML是经过以市场组合作为切点组合的一条资本配置线。CML的等式为

$$E(R_p) = R_F + \frac{E(R_M) - R_F}{\sigma_M} \sigma_p \quad (11-8)$$

式中, $E(R_p)$ 表示在资本市场线上的投资组合 p 的期望收益率;

R_F 表示无风险利率;

$E(R_M)$ 表示市场组合的期望收益率;

σ_M 表示市场组合收益率的标准差;

σ_p 表示投资组合 p 收益率的标准差。

CML的斜率, $\frac{E(R_M) - R_F}{\sigma_M}$ 被称为风险的市场价格(market price of risk),这是因为它反映了每单位市场风险所要求的市场风险溢价。正如我们所指出的,CML只描述了有效投资组合的期望收益率。资本市场线的含义是所有均值一方差投资者,无论他们的风险容忍度是怎样的,都会通过将无风险资产与所有风险资产的市场组合(这一单个风险资产)进行组合来满足其投资需求。

在下一节中,我们将给出一个描述任一资产或者投资组合(无论是有效还是无效)的均值一方差理论。

11.2.6 资本资产定价模型

资本资产定价模型(capital asset pricing model, CAPM)自从20世纪60年代初提出以来,在定量投资管理的发展过程中发挥着关键性的作用。在本节中,我们将回顾一些有关它的重要内容。

CAPM做出了如下假设:^①

- 投资者只需要知道期望收益率、收益率的方差和协方差,就能确定出哪个投资组合对他们来说是最优的。(该假设出现在所有的均值一方差理论之中。)
- 投资者对于风险资产的平均收益率,收益率的方差和相关系数具有相同的看法。
- 投资者可以以任何数量买卖资产而不会对其价格产生影响,并且所有的资产都是有市场的(即所有资产都可以随时进行交易)。
- 投资者可以以无风险利率进行借贷(没有数量上的限制),而且可以对任何资产以任何数量进行卖空。

^① 完整的假设列表,参见 Elton, Gruber, Brown 和 Goetzmann (2003)。

- 投资者对于收益不需要支付税收,而且对于交易无需支付交易费。

当 CAPM 的假设成立时, CML 就代表有效前沿。因此, 在 CAPM 的世界里, 所有投资者可以通过将无风险资产和相同的切点组合 (其是所有风险资产 (没有风险资产被排除在外) 的市场组合) 进行组合来满足其投资需求。

下面的等式描述了所有资产和投资组合 (无论有效还是无效) 的期望收益率:

$$E(R_i) = R_F + \beta_i [E(R_M) - R_F] \quad (11-9)$$

式中, $E(R_i)$ 表示资产 i 的期望收益率;

R_F 表示无风险收益率;

$E(R_M)$ 表示市场组合的期望收益率。

$$\beta_i = \text{Cov}(R_i, R_M) / \text{Var}(R_M)$$

式 (11-9) 本身被称为资本资产定价模型 (capital asset pricing model), 它的图像被称为证券市场线 (security market line, SML)。CAPM 是一个将任何资产 (或者投资组合) 期望收益率描述为关于其贝塔 (beta) β_i (它是度量资产对于市场变动敏感性的一个指标) 线性函数的一个等式。CAPM 认为期望收益率由两部分组成: 首先是无风险利率, 其次是等于 $\beta_i [E(R_M) - R_F]$ 的额外收益率。 $[E(R_M) - R_F]$ 这项是预期超出市场的超额收益率。这个量是市场风险溢价 (market risk premium); 如果我们 100% 投资于市场组合, 那么市场溢价就是我们平均所期望获得的相对于无风险利率的额外收益率。

市场风险溢价会与资产的贝塔相乘。贝塔为 1 表示平均市场的敏感性, 我们认为具有该贝塔的资产会完全获得市场风险溢价。^① 比 1 大的贝塔表示大于平均市场的风险, 根据 CAPM, 它将获得更高的预期超额收益率。相反, 低于 1 的贝塔表示低于平均市场的风险, 根据 CAPM, 它将获得更低的预期超额收益率。预期超额收益率只与表示为贝塔的市场风险有关。市场收益率的敏感性是造成资产间不同预期超额收益率的唯一原因。^②

像所有基于理论的模型一样, CAPM 来自一组假设。CAPM 描述了在假设条件下的一个金融市场的均衡, 如果模型是正确的, 而任一资产的期望收益率不同于由 CAPM 给出的期望收益率, 那么在这个时候, 市场的力量将会发挥作用。该力量将促使其收益率回复到由该模型所设定的关系。例如, 一个提供比其贝塔所得到的期望收益率更高的股票, 它的股价将会升高, 从而降低该股票的预期收益率; 投资者将认为一个投资广泛的投资组合将抵消任何该股票所含有的非市场风险。

因为由 CAPM 所定义的市场组合包含了所有的资产, 所以其是不可观测的。在实践中, 我们必须用一些包含广泛资产的指数来代表它。由于 CAPM 已经被主要用来对股票进行估值, 因此, 对于市场组合的一个常用的选择是一个广泛投资的市值加权的指数或者市场的一个代理变量。期望收益率和贝塔间的直线关系源于市场组合的有效性。所以, CAPM 理论等价于认为不可观测的市场组合是有效的, 但是这并不意味着任何市场的代理变量都是有效的。^③ 对于金融从业者来说,

① 市场组合本身的贝塔为 1, 因为 $\beta_M = \text{Cov}(R_M, R_M) / \text{Var}(R_M) = 1$ 。因为市场组合包含了所有的资产, 所以平均持有的资产一定具有 1 的贝塔。如果我们使用代表市场的指数来计算一个指数中资产的贝塔, 那么我们会采用与上面相同的结论, 即假设市场指数的贝塔为 1。

② 这个观点的一个直观想法是市场是一个完全分散化的投资组合。我们可以通过持有市场组合来排除任何其他的风险, 并且我们可以无成本地持有市场组合 (根据无交易成本假设)。即使关于个人资产 (诸如人力资本, 表示为收入能力) 的风险也可以被充分分散化 (所有资产都是可交易的)。投资者不应该对于他们可以无成本对冲的风险要求额外的收益率。

③ 更多这方面的内容, 参见 Bodie, Kane 和 Marcus (2001)。

比起 CAPM 作为一个理论是否是一个完全的真理,他们更关心由可获得的市场代理变量所算得的贝塔是否对于评估不同投资策略的期望平均收益率有用。而目前的证据则倾向于认为存在多种影响投资策略平均收益率的系统性风险。

11.2.7 均值方差投资组合选择规则:一个介绍

在这一节中,我们将介绍从均值一方差的角度所得到的一些投资组合选择的原则。最基本的投资组合选择决策之一是从一组可选择的资产类中选出一个最优的资产配置。而另一类决策涉及对于某个已知投资组合的改进。因为我们在得到一个比已知投资组合在均值一方差上更优的投资组合时,无需确认改进后的该组合是否是最优的,所以这种决策方法更为简单。我们下面将从第2种决策类型的简要讨论中开始我们的介绍。

11.2.7.1 与已知投资组合相关的决策

我们考察与已知投资组合相关的2种决策,在决策的过程中均值一方差分析将发挥重要作用。

独立投资的投资组合之间的比较

马克维茨决策规则 (Markowitz decision rule) 提供了这样的—个准则,均值一方差投资者有时利用该准则可以在面对将其资金全部投资到资产 A 或者资产 B 的选择决策时做出判断。当以下2个条件之一成立时,该投资者偏好 A 而不是 B:

- A 的平均收益率大于等于 B,但是 A 的收益率具有比 B 更小的标准差。
- A 的平均收益率严格大于 B,但是 A 和 B 的收益率具有相同的标准差。

当根据马克维茨决策规则判断, A 优于 B 时,我们就说 A 在均值一方差上优于 B: 资产 A 显然比 B 更有效地利用了风险。例如,如果一个投资者被要求在 (1) 平均收益率为 9%, 标准差为 12% 的资产配置 A 和 (2) 平均收益率为 8%, 标准差为 15% 的资产配置 B 之间进行选择,那么该均值一方差投资者将更喜欢备选 A, 因为它预期能提供更高的平均收益率和更小的风险。需要注意的一点是,当配置资产同时具有更高的平均收益率和标准差时,马克维茨决策规则将无法选择出哪个资产配置更优;相反,该选择将依赖于个人投资者的风险容忍程度。

如果我们可以以无风险利率借贷,那么我们就找到一组扩展的均值一方差占优投资组合间的关系。接着,我们可以使用无风险资产来匹配用来进行比较的不同投资组合的风险。这时,夏普比率 (平均收益率超出无风险收益率部分与收益率标准差的比率) 将成为一个恰当的度量指标。

如果一个投资组合 p 比另一个投资组合 q 具有更高的正的夏普比率 (如果以无风险利率借贷是可能的话) 那么 p 就在均值一方差上优于 q 。[⊖]

假如资产配置 A 和之前假设的一样,但是现在 B 的平均收益率为 6%, 标准差为 10%。配置 A 的平均收益率高于 B (9% 对 6%), 但是风险也更高 (12% 对 10%), 此时马克维茨决策规则将无法得出那个配置更好的结论。假如我们可以以 3% 的无风险收益率进行借贷。A 的夏普比率 $(9 - 3)/12 = 0.50$, 高于 B 的夏普比率 $(6 - 3)/10 = 0.30$, 所以我们能得到 A 在均值一方差上优于 B。注意到 83.3% 投资于 A, 16.7% 投资于无风险资产的投资组合具有和 B 一样的标准差,因

⊖ 该命题相反的结论也是正确的: 如果一个投资组合 p 在均值一方差上优于另一个投资组合 q , 那么 p 比 q 具有更高的正的夏普比率。该命题的证明可以在 Dybvig 和 Ross (1985a) 中找到。注意, 这里我们对于更高夏普比率的投资组合假设了一个正的夏普比率的条件, 由此排除了一些在负夏普比率投资组合间进行比较时产生的反例。

为 $0.833(12) = 10\%$ ，平均收益率为 $0.833(9) + 0.167(3) = 8\%$ ，对于 B 的 6%。总之，我们将更高夏普比率的投资组合 p 与无风险资产进行组合后得到与投资组合 q 具有相同风险，但其平均收益率更高的投资组合。根据 B 初始的定义（平均收益率为 8%，而标准差为 15%），B 的夏普比率为 0.33，所以正如我们所预计的，根据夏普比率所做出的决策和根据马克维茨决策规则所得到的结果是一致的。

在实践中，当我们对于充分分散化的投资组合进行相互比较并且所选择的投资组合收益率的分布至少是近似正态的时候，那么上面决策方法的可信度是很高的。

在已知投资组合中加入一项投资的决策

我们已经将在 2 个资产配置间进行选择的方法描述为一个 2 选 1 的命题。我们现在要讨论是否要在已有投资组合中加入新资产类，或者说进一步分散化已有投资组合的问题。

假如你持有一个投资组合 p ，其期望或者平均收益率为 $E(R_p)$ ，收益率的标准差为 σ_p 。那么你被给予在你的资产组合中增加另一个投资的机会，比如增加一个资产类。在你的投资组合中包含一项新投资的正头寸是否会对均值一方差有所改善？为了回答这个问题，你需要 3 个输入变量：

- 新投资的夏普比率
- 已有投资组合的夏普比率，以及
- 新投资收益率与投资组合 p 收益率间的相关系数， $\text{Corr}(R_{\text{new}}, R_p)$

如果下面的条件成立，那么在你的投资组合中增加新的资产是最优的：[⊖]

$$\frac{E(R_{\text{new}}) - R_F}{\sigma_{\text{new}}} > \left(\frac{E(R_p) - R_F}{\sigma_p} \right) \text{Corr}(R_{\text{new}}, R_p) \quad (11-10)$$

该表达式说明为了通过增加新投资的持有来获得收益，新投资的夏普比率必须大于你已有投资组合的夏普比率与新投资收益率和现有投资组合收益率间相关系数的乘积。如果式 (11-10) 成立，我们就可以将新投资与之前的投资进行组合以得到一个更优的风险资产有效前沿（其中的切点组合具有更高的夏普比率）。注意，虽然表达式可以给出我们能否通过增加一项新资产的正头寸来在边际上对均值一方差进行改善，但是它并没有告诉我们应该增加持有多少新资产的投资，或者更一般地，包含新资产的有效前沿可能是怎样的——为了确定该有效前沿，我们将需要求解一个最优化问题。由式 (11-10) 我们可以看到，与我们将投资作为独立情形而只需要考虑夏普比率（从一个均值一方差的视角）来进行选择的例子不同，在这个例子中我们能够组合 2 个投资，并且我们还必须考虑它们的相关系数。例 11-6 说明了如何使用式 (11-10) 来确定是否要增加一个资产类。

例 11-6 增加一种资产类的决策

吉姆·瑞格 (Jim Regal) 是加拿大一家投资于加拿大股票、债券、房地产和美国股票的养老基金的投资总监。投资组合的夏普比率为 0.25。投资委员会考虑增加一个或者其他另一个（但是不会同时增加 2 个）下列的资产类：

- 欧洲债券：预计的夏普比率为 0.10；预计与现有投资组合的相关系数为 0.42。
- 非北美发达市场股票，以 MSCI EAFE（欧洲 (Europe)，澳大利亚 (Australasia)，远东 (Far East)）指数表示：预计的夏普比率为 0.30；预计与现有投资组合的相关系数为 0.67。

⊖ 参见 Blume (1984) 和 Elton, Gruber, 和 Rentzler (1987)

(1) 解释投资委员会是否应该将欧洲债券增加到已有投资组合之中。

(2) 解释投资委员会是否应该将非北美发达市场股票增加到已有投资组合之中。

解:

(1) (已有投资组合的夏普比率) $\times$ (欧洲债券与已有投资组合的相关系数) $= 0.25(0.42) = 0.105$ 。如果欧洲债券的预期夏普比率超过 0.105, 那么我们就应该增加欧洲债券。因为投资委员会预计欧洲债券的夏普比率为 0.10, 因此, 委员会不应该将欧洲债券增加到已有投资组合之中。

(2) (已有投资组合的夏普比率) $\times$ (新股票类与已有投资组合的相关系数) $= 0.25(0.67) = 0.1675$ 。因为预计的非北美发达市场股票的夏普比率超过 0.1675, 所以投资委员会应该将它们增加到已有投资组合之中。

在例 11-6 中, 即使养老基金已有投资组合与推荐的新股票类的相关系数为 +1, 此时增加新资产类也不会带来潜在的风险降低的好处, 而式 (11-10) 也表明应该增加该资产类, 这是因为增加资产类的条件得到了满足, 即 $0.30 > 0.25(0.1)$ 。对于任一投资组合, 我们总是可以通过增加一个比原有投资组合更高夏普比率的投资来从边际上改善均值方差情况。该结果是符合直观想象的, 因为更高夏普比率的投资将在均值一方差上优于已有投资组合。这里, 我们要再次强调要使这些结论成立, 均值一方差分析的假设必须得到满足。

11.2.7.2 确定一个资产配置

在这一节中我们的目标是总结关于资产配置均值一方差的一些观点。在先前的章节中, 我们在一个最简单的例子中, 从数学上给出了确定一个风险资产最小方差前沿的目标和约束条件, 在这个例子中唯一对于投资组合权重的约束是它们加总必须为 1。利用式 (11-3) 加上反映不允许卖空的 $\omega_i \geq 0$ 约束来确定有效前沿是许多机构投资者在确定一个资产配置时的出发点。^①于是, 如果投资者可以以无风险利率进行借贷, 那么均值一方差理论指出所选择的资产配置可以表示为切点组合。投资经理就会将切点组合与无风险资产进行组合以达到其期望风险水平下的有效投资组合。因为切点组合代表了最高夏普比率的风险资产的投资组合, 所以将这个准则作为资产配置的一个基准是合理的。

然而, 这一理论将把无风险资产的投资处理为一个容易获得且经过风险调整的变量。而在实际投资中, 投资者可能会被限制使用保证金 (借款), 他们可能会面对风险资产正头寸最大和最小金额的限制, 或者他们可能具有其他无法满足均值一方差理论假设的原因。^②在这种情况下, 投资者可能会在不同于切点组合的无风险资产类中建立一个资产配置。以收益率的标准差来量化其风险容忍程度, 该投资者可以在风险资产的有效前沿上选择与其选择的标准差水平相对应的投资组合。

CAPM 框架甚至提供了更狭小的选择空间, 因为它增加了投资者对于平均收益率、收益率方差和相关系数具有相同看法的假设。于是所有的投资者将选择相同的切点组合, 即以所有风险资产市值加权的市场组合。该投资组合代表最大可能的风险分散化水平。如此一致的包含所有资产

① 参见 Grinblatt 和 Titman (1998) 以及 Elton, Gruber, Brown 和 Goetzmann (2003) 中关于求解方法和机制的讨论。在实务中, 投资者经常会增加其他的约束, 例如最大百分比的约束, 以获得一个可信的结果。我们将在最小方差前沿的不稳定性一节中更深入地讨论这一方法。

② 例如, 如果分析是基于实际利率考虑的话, 那么无风险资产可能不容易获得。只有短期抗通胀证券, 如果可以获得的话, 才是在这种情况下真正的无风险资产。

的投资组合一般当然是无法得到的。然而，在实际中，投资者能够持有投资于大部分全球资产类的高度分散化的被动管理投资组合，其近似反映所包含资产的相对市场价值。该资产配置可以用来解释人们不同的预期。例如，布莱克-利特曼（1991，1992）资产配置模型采用了一个完全分散化的市值加权的资产配置作为投资者的一个中立的出发点。该模型包含了将市值权重调整到反映投资者对于不同于均衡模型（CAPM）关于期望收益率观点的步骤。

均值一方差理论对于投资组合建立和资产配置的内容已经进行了深入的讨论，而研究者们也已经发现了它的许多缺陷。我们将在之后的一节中讨论一个与最优化步骤敏感性相关的最重要问题——有效前沿的不稳定性。我们必须认识到作为一个单期模型，均值一方差分析忽视了流动性和税收的考虑，而这在投资者会重新调整投资组合权重的多期情况下是会产生重要影响的。与此类似地，与最优化过程相关的时间长短，通常是一年，也可能会因为长时间段得到输入变量的困难，而使得实际求解的区间比投资者实际的时间区间更短。^①均值一方差分析只考虑不同资产类收益率在单期内之间的相关系数，但不是某个给定资产收益率的序列相关性（长期和短期的依赖性）。蒙特卡罗模拟有时被用于资产配置，来解释这类多期问题。^②尽管该理论有很多缺陷，但是均值一方差分析仍在选择一个资产配置中提供了我们一个目标以及相对可接受的从所面对的无限的选择中缩小可投资资产数量的一种方法。^③

11.3 均值方差分析在应用中的问题

我们现在讨论在应用均值一方差分析选择投资组合过程中可能产生的一些实际问题。我们关注的2个问题是

- 均值一方差最优化问题的输入变量（inputs）的估计；
- 最小方差前沿的不确定性，这个问题是由最优化过程对于输入变量的敏感性所造成的。

对于第1个问题，我们会存在2个关于期望收益率、方差和相关系数的主要问题。首先，哪种估计方法是可行的？其次，哪种估计方法是最准确的？而对于最优化过程的敏感性问题，我们需要询问这样2个问题，首先，该问题的原因是什么？其次可以解决它的补救方法是什么？

11.3.1 估计均值方差优化问题中的输入参数

在本节中，我们将比较计算均值一方差最优化输入变量的不同方法的可行性和准确性。这些方法会使用下列变量中的其中一组来对于参数进行估计：

- 历史均值、方差和相关系数
- 利用市场模型估计的历史贝塔
- 调整的贝塔

11.3.1.1 历史估计值

该方法涉及直接从历史数据来计算均值、方差和相关系数。即使我们对于中等数目的资产进

① 参见 Swenson（2000）。例如，如果投资者的时间区间是5年，那么形成一个有效前沿涉及估计5年收益率的相关系数。许多资产具有有限独立的5年观测值数据，从而使得我们难以估计所需的相关系数。

② 参见概率分布一章中更多的关于蒙特卡罗模拟的内容。

③ 例如，Chow（1995）采用均值一方差最优化方法来回答投资经理关于相对于基准的业绩表现的看法，而 Chow，Jacquier，Kritzman 和 Lowry（1999）采用最优化方法解释了在经济萧条的时候相关系数可能发生改变的情况。

行最优化,历史方法也要求估计很多参数。因此,进行资产配置比形成一个包含许多股票的投资组合更加现实、可行。

一位投资组合经理需要估计的确定最小方差的参数数量依赖于投资组合中潜在的股票数目。如果一位投资组合经理的投资组合中有 n 只股票,并希望使用均值一方差分析,她必须估计:

- n 只股票的期望收益率的参数;
- n 只股票收益率方差的参数;
- $n(n-1)/2$ 个所有股票收益率之间的协方差参数;
- 总共有 $n^2/2 + 3n/2$ 个参数。

历史方法的两大缺陷分别是需要估计的参数数目庞大以及历史估计出的输入变量的准确性问题。

所需估计参数的数目轻易地就会变得很大,这主要是因为协方差的数目以证券数目的平方的速度而增长。如果投资组合经理想要计算 100 只股票投资组合的最小方差前沿,那么她将需要估计 $100^2/2 + 3(100)/2 = 5150$ 个参数。如果她想要计算 1000 只股票投资组合的最小方差前沿,那么她将需要估计 501500 个参数。这项工作不仅非常无趣,而且这项工作几乎是不可能完成的。^①

第2个缺陷是收益率参数的历史估计值一般具有很大的估计误差。该问题对于方差的估计值不会产生严重的影响。^②但是该问题对于平均收益率的历史估计值很敏感,因为风险资产收益率相对于均值的变动会很大,并且该问题不能通过增加观测次数而得到改善。估计误差同样也在协方差历史估计值中非常严重。协方差估计中存在问题的原因是历史方法本质上是试图刻画每个历史数据集的随机特征,从而降低了估计值在预测方面的作用。在对于 1973~1997 年美国股票月度收益率的研究中,陈(Chan)、卡塞斯基(Karceski)和兰考尼肖科(Lakonishok)(1999)发现 36 个月区间的过去和未来样本协方差间的相关系数为 0.34,但是在 12 个月区间的相关系数只有 0.18。

在当前业界的实践中,历史的样本协方差矩阵在未调整之前不会被使用在均值一方差最优化问题中。调整后的方差和协方差的值可能分别是原始样本值和方差平均值或者协方差的加权平均值。例如,如果一只股票月度收益率的方差为 0.0210 而平均股票月度收益率方差为 0.0098,那么该方法可能将 0.0210 向下方均值的方向进行调整。将数值向着均值的方向调整能降低由样本误差所造成的估计值偏差。^③

在估计平均收益率时,分析师使用了许多不同的方法。他们可以调整历史平均收益率来反映当前市场条件和历史平均经验之间的反差。他们也可以不断地使用估值模型,如基于预期未来现金流的模型或者均衡模型如 CAPM 来建立未来平均收益率的估计值。他们使用的这些方法不仅反映了估计误差的技术问题,而且体现了在未来业绩将由过去平均业绩所反映的假设下的风险。

11.3.1.2 市场模型的估计值:历史贝塔(未经过调整)

一种计算资产收益率方差和协方差更简单的方法涉及这样的观察结果:资产收益率可以通过

① 时间序列观测值的数量必须超过证券的数量,以使得协方差(包括方差)得以估计。

② 参见 Chan, Karceski 和 Lakonishok (1999) 中关于未来方差比未来的协方差相对更能通过历史数据进行预测的实证证据。

③ 关于该方法(即被称为缩减估计量的方法)的更多内容,参见 Michaud (1998) 以及 Ledoit 和 Wolf (2004)。

由有限的一组变量或者因素的相关系数来反映出它们彼此之间的相关关系。这类模型中最简单的模型就是市场模型，它描述了一个资产收益率与市场组合收益率间的回归关系。对于资产 i ，该资产的收益率可以用模型描述为

$$R_i = \alpha_i + \beta_i R_M + \epsilon_i \quad (11-11)$$

式中， R_i 表示资产 i 的收益率； R_M 表示市场组合的收益率； α_i 表示资产 i 平均收益率与市场收益率不相关的部分； β_i 表示资产 i 收益率对于市场组合收益率的敏感性； ϵ_i 表示误差项。

我们首先考虑如何解释 β_i 。如果市场收益率增长一个百分点，那么市场模型预测资产 i 的收益率将增长 β_i 个百分点。（记住 β_i 是市场模型的斜率。）

现在考虑如何来解释 α_i 。如果市场收益率为 0，那么市场模型预测资产 i 的收益率将是 α_i ，即该市场模型的截距。

该市场模型对于式 (11-11) 中的项做出如下这些假设：

- 误差项的期望值为 0，因此 $E(\epsilon_i) = 0$ 。
- 市场收益率 (R_M) 与误差项不相关， $\text{Cov}(R_M, \epsilon_i) = 0$ 。
- 不同资产间的误差项 ϵ_i 都不相关。例如，资产 i 的误差项与资产 j 的误差项是不相关的。

因此， $E(\epsilon_i \epsilon_j) = 0$ ，对于所有不相等的 i 和 j 。[⊖]

注意到这些假设中的一些内容与我们在相关性和回归一章中对于单变量线性回归模型所做的假设是非常相似的。然而，市场模型并没有假设误差项服从正态分布以及资产间误差项的方差是相同的。

给定这些假设，市场模型做出如下对于资产期望收益率和资产收益率方差和协方差的预测。[⊖]

首先，资产 i 的期望收益率依赖于市场的期望收益率 $E(R_M)$ ，资产 i 的收益率对于市场收益率的敏感性 β_i ，以及资产 i 的收益率与市场收益率独立的部分 α_i 。

$$E(R_i) = \alpha_i + \beta_i E(R_M) \quad (11-12)$$

其次，资产 i 收益率的方差依赖于市场收益率的方差 σ_M^2 ，市场模型中资产 i 收益率误差的方差 $\sigma_{\epsilon_i}^2$ 以及敏感性 β_i 。

$$\text{Var}(R_i) = \beta_i^2 \sigma_M^2 + \sigma_{\epsilon_i}^2 \quad (11-13)$$

在该模型中市场组合是风险的唯一来源，式 (11-13) 的第 1 项 $\beta_i^2 \sigma_M^2$ 有时被称为资产 i 的系统性风险。式 (11-13) 中的误差方差项 $\sigma_{\epsilon_i}^2$ 有时被称为资产 i 的非系统性风险。

第三，资产 i 的收益率与资产 j 的收益率间的协方差依赖于市场收益率的方差 σ_M^2 以及敏感性 β_i 和 β_j 。

$$\text{Cov}(R_i, R_j) = \beta_i \beta_j \sigma_M^2 \quad (11-14)$$

我们可以使用市场模型来极大地降低得到均值一方差最优化所需输入变量的计算规模。对于 n 个资产的每一个，我们只需要知道 α_i ， β_i ， $\sigma_{\epsilon_i}^2$ 以及市场期望收益率和方差。因为我们在市场模型中只需要估计 $3n + 2$ 个参数，所以与估计资产收益率的历史均值、方差和协方差相比，我们只需要很少的参数来构建最小方差前沿。例如，如果我们估计 1 000 个资产（比如说 1 000 只不同股

⊖ $\text{Cov}(\epsilon_i, \epsilon_j) = E\{[\epsilon_i - E(\epsilon_i)][\epsilon_j - E(\epsilon_j)]\} = E\{(\epsilon_i - 0)(\epsilon_j - 0)\} = E(\epsilon_i \epsilon_j) = 0$ 。不相关误差的假设不是无害的。如果我们具有超过一个影响资产收益率的因素，那么该假设将是不正确的，单变量模型将得到不准确的资产收益率协方差的估计值。

⊖ 参见 Elton, Gruber, Brown 和 Goetzmann (2003) 中关于这些结果的推导。

票)的最小方差前沿,市场模型将使用3 002个参数来构建最小方差前沿,而历史估计方法,正如之前所讨论的,将需要501 500个参数。

由于我们不知道市场模型的参数,因此我们必须估计它们。但是我们应该采取何种方法呢?最简便的方法是利用关于市场收益率和每种资产收益率的时间序列数据来估计一个线性回归。

我们可以通过对于每个资产分别采用线性回归,并使用资产收益率和市场收益率的历史数据来用市场模型估计 α_i 和 β_i 。^①该回归的结果给出了 β_i 的估计值 $\hat{\beta}_i$;我们称该估计值为未调整的贝塔。之后我们将介绍调整后的贝塔。我们可以使用这些估计量来计算这些均值—方差最优化问题所需的收益率的均值、方差和协方差。

例 11-7 使用市场模型计算股票的相关系数

你正在估计思科系统 (Cisco Systems) (纳斯达克代码: CSCO) 和微软 (Microsoft) (纳斯达克代码: MSFT) 在 2003 年年末收益率之间的相关系数。你基于这 2 只股票的月度收益率,对于市场模型进行回归,其中的市场组合采用的是标准普尔 500 指数。你获得下列回归结果:

思科公司估计出的贝塔 $\hat{\beta}_{\text{CSCO}}$ 为 2.09,而残差的标准差为 $\hat{\sigma}_{\epsilon \text{CSCO}}$ 为 11.52。

微软公司估计出的贝塔 $\hat{\beta}_{\text{MSFT}}$ 为 1.75,而残差的标准差为 $\hat{\sigma}_{\epsilon \text{MSFT}}$ 为 11.26。

你对于标准普尔 500 月度收益率的方差估计为 $\sigma_M^2 = 29.8$,其对应的年度收益率标准差约为 18.9%。使用给定的数据,估计思科公司和微软公司收益率的相关系数。

解:

我们计算出 $\sigma_{\epsilon \text{CSCO}}^2 = 132.71$ 和 $\sigma_{\epsilon \text{MSFT}}^2 = 126.79$ 。使用相关系数的定义,即协方差除以每个变量的标准差,再利用式 (11-13) 和式 (11-14),我们有

$$\begin{aligned} & \frac{\text{Cov}(R_{\text{CSCO}}, R_{\text{MSFT}})}{\text{Var}(R_{\text{CSCO}})^{1/2} \text{Var}(R_{\text{MSFT}})^{1/2}} \\ &= \frac{\hat{\beta}_{\text{CSCO}} \hat{\beta}_{\text{MSFT}} (\hat{\sigma}_M^2)}{[\hat{\beta}_{\text{CSCO}}^2 (\hat{\sigma}_M^2) + \hat{\sigma}_{\epsilon \text{CSCO}}^2]^{1/2} [\hat{\beta}_{\text{MSFT}}^2 (\hat{\sigma}_M^2) + \hat{\sigma}_{\epsilon \text{MSFT}}^2]^{1/2}} \\ &= \frac{(2.09)(1.75)(29.8)}{[(2.09)^2(29.8) + 132.71]^{1/2} [(1.75)^2(29.8) + 126.79]^{1/2}} = 0.4552 \end{aligned}$$

因此,市场模型预测这 2 个资产收益率间的相关系数为 0.4552。

使用市场模型的一个困难之处是确定出一个代表市场的恰当指数。一般来说,使用市场模型来确定单个国内股票风险的分析师会使用国内股票市场的指数收益率来代表市场。在美国,该指数可能是标准普尔 500 或者威尔谢尔 5000 指数;在英国,可能会采用金融时报股票交易所 100 指数。使用股票市场指数的收益率可能会对股票给出一个合理的市场模型,但是它对于其他资产类可能并不能给出合理的模型。^②

① 一种常用的方法是使用 60 个月度收益率来估计该模型。彭博 (Bloomberg) 终端的默认设定是采用 2 年的周数据来估计该模型。

② 使用该模型来估计其他资产类的风险可能违背之前所讨论的单因素模型的 2 个假设: 市场收益率 R_M 与误差项 ϵ_i 独立; 以及不同资产间的误差项 ϵ_i 独立。如果这 2 个假设的任一个被违背,那么市场模型将无法给出准确的期望收益率以及收益率方差和协方差的估计值。

11.3.1.3 市场模型的估计值：调整后的贝塔

我们可以使用由市场模型所估计出的历史贝塔来进行均值一方差最优化吗？在我们回答这个问题之前，我们需要重申一下我们的目标：我们想要预测一组资产的期望收益率以及这些收益率的方差和协方差，以使得我们可以估计出这些资产的最小方差前沿。基于历史贝塔的估计值依赖于这样一个关键假设：某一资产的历史贝塔是该资产未来贝塔的最优估计值。如果贝塔随时间发生变化，那么该假设是不正确的。因此，我们可能想要使用一些其他的度量指标来代替历史贝塔来对于该资产未来的贝塔进行估计。这些其他的预测值一般被称为调整后的贝塔（adjusted beta）。研究者们已经证明调整后的贝塔通常能比历史贝塔给出更好的未来贝塔的预测值。因此，从业者们经常采用调整后的贝塔。

例如，假设我们在第 t 期，我们想要估计一组股票在 $t+1$ 期的最小方差前沿。我们需要使用在第 t 期可以获得的数据来预测 $t+1$ 期股票的期望收益率、方差以及协方差。然而，注意到在第 t 期对于某只股票所估计出的历史贝塔可能不是我们在第 t 期对于 $t+1$ 该股票未来贝塔的最优估计值。而 $t+1$ 期的最小方差前沿必须基于 $t+1$ 期的贝塔预测值。

如果每只股票的贝塔从一期到另一期都服从一个随机游走，那么我们可以将第 t 期股票 i 的贝塔和第 $t+1$ 期股票 i 的贝塔间的关系表示为

$$\beta_{i,t+1} = \beta_{i,t} + \epsilon_{i,t+1}$$

式中， $\epsilon_{i,t+1}$ 是误差项。如果贝塔服从随机游走，那么 $\beta_{i,t+1}$ 的最佳预测值为 $\beta_{i,t}$ ，这是因为误差项的均值为 0。历史贝塔将是未来贝塔的最佳预测值，而历史贝塔将不需要调整。

在实际中，每只股票的贝塔通常在一期与下一期之间并不服从随机游走，因此，历史贝塔并不一定是未来贝塔的最佳预测值。例如，如果贝塔可以表示成如下一阶自回归的形式，那么

$$\beta_{i,t+1} = \alpha_0 + \alpha_1 \beta_{i,t} + \epsilon_{i,t+1} \quad (11-15)$$

如果我们使用历史贝塔的时间序列数据估计式 (11-15)，那么 $\beta_{i,t+1}$ 的最佳预测值将为 $\hat{\alpha}_0 + \hat{\alpha}_1 \beta_{i,t}$ ，而不是 $\beta_{i,t}$ 。

用调整后的贝塔预测未来贝塔要比用历史贝塔预测更好，这是因为贝塔一般是均值回复的。[⊖]因此，我们应该使用调整后的，而不是历史的贝塔。一种从业者常用的调整贝塔的方法是假设 $\alpha_0 = 0.333$ 而 $\alpha_1 = 0.667$ 。在这种调整下，

- 如果历史贝塔等于 1.0，那么调整后的贝塔将为 $0.333 + 0.667(1) = 1.0$ 。
- 如果历史贝塔等于 1.5，那么调整后的贝塔将为 $0.333 + 0.667(1.5) = 1.333$ 。
- 如果历史贝塔等于 0.5，那么调整后的贝塔将为 $0.333 + 0.667(0.5) = 0.667$ 。

因此，贝塔的均值回复水平为 1.0。如果历史的贝塔高于 1.0，那么调整后的贝塔将低于历史的贝塔；如果历史的贝塔低于 1.0，那么调整后的贝塔将高于历史的贝塔。[⊖]

11.3.2 最小方差前沿的不稳定性

虽然如式 (11-3) 所表示的标准的均值一方差最优化理论对于构建投资组合是一个方便而且

⊖ 例如，参见 Klemkosky 和 Martin (1975)。

⊖ 虽然从业者们通常使用这种方法来计算调整后的贝塔，但是我们没有看到任何出版的研究表明 $\alpha_0 = 0.333$ 和 $\alpha_1 = 0.667$ 是在计算调整后的贝塔时所应该使用的最佳的系数。基本面贝塔（fundamental betas）基于公司的基本面数据（市盈率，盈利增长率，市值，波动率等）来预测贝塔。咨询公司诸如 BARRA 销售这些基本面贝塔的估计值。

客观的方法，但是我们在实践中对于其结果进行解释时必须加以小心。在本节中，我们将讨论关于均值一方差最优化理论在使用中所必须注意的地方。该方法可能产生的问题已经得到了人们广泛地研究，而对于这些问题的补救方法也已经被人们所提出。在了解了这些知识之后，均值一方差最优化方法将仍旧是一个有用的工具。

均值一方差最优化方法的一个主要问题是输入变量的微小变动将引起最小方差前沿（有效前沿）的巨大变动。该问题被称为最小方差前沿的不稳定性（instability in the minimum-variance frontier）。该问题的产生是因为在求解最小方差前沿过程中所使用的期望收益率、方差和协方差的不确定性所造成的。

例如，假设我们使用历史数据来计算最优化问题中所使用的估计值。这些均值、方差和协方差都是受到随机变动影响的样本估计值。而在抽样一章中，我们讨论了样本均值是如何具有一个概率分布的，该分布被称为抽样分布。样本均值只是相应总体均值的点估计值。^①最优化过程试图最大地利用这些资产估计值间的差别。当这些差别在统计上（和经济上）不显著（例如表示为随机变动）时，所得到的最小方差前沿就会具有误导性，而且其在实践上也是没有用处的。因此，均值一方差最优化方法过于强调对于数据进行拟合，而其对于实际上并没有意义的资产估计值间的差别过分地予以关注。在没有卖空限制的最优化问题中，所求解出的资产可能会被赋予很严重的负权重（赋予一个资产负的权重意味着该资产被卖空），而这正反映了其过度地进行了拟合。具有很大卖空头寸的投资组合在实践中是没有任何意义。^②由于对于输入变量微小变化的敏感性，均值一方差最优化问题可能会建议人们非常频繁地对于投资组合进行再平衡，而这会带来很高的成本。而对于不稳定性问题的解决方法，主要包括如下这些：

- 增加不能卖空的约束（有时也会增加机构投资要求）。在式（11-3）中，我们可以增加一个不能卖空的约束，即规定所有资产权重必须为正： $\omega_j \geq 0, j = 1, 2, 3, \dots, n$ 。^③
- 改善最优化问题中输入变量的统计性质。
- 使用有效前沿的统计概念，其反映了最优化问题的输入变量是随机变量而不是常数的情况^④。

我们上面论述了均值一方差最优化理论可能会给出过于频繁的投资组合再平衡。类似地，当利用不同时间段的历史数据进行计算时，我们发现最小方差前沿一般是不稳定的。一个可能的解释是不同的前沿反映了样本时间段之间资产收益率分布的变化。最小方差前沿在时间上的不稳定性也可能源于均值、方差和协方差的随机变动，而相应参数实际上并没有发生变化。用于均值一方差最优化问题的样本时间段的微小不同可能极大地影响一个模型，即使资产收益率的分布是平稳的。例11-8给出了用于最优化问题的数据在时间上的不稳定性。

① 相应资产收益率的均值非常难以准确地估计。参见 Luenberger（1998）关于该问题的介绍，或者参见 Black（1993）。

② 在实践中，很少有从事卖空的投资者会因为对于均值、方差和相关系数的分析而持有很大的卖空头寸。一个卖空头寸的损失可能是无限的。

③ 在实践中，其他专门关于头寸大小的限制条件有时也被使用。

④ 例如，Michaud（1998）定义了给定置信水平下统计上等价的有效投资组合的范围。落入该范围的投资组合与有效组合是一致的，从而不需要进行再平衡。

例 11-8 最小方差前沿在时间上的不稳定性

在例 11-2 中，我们计算出了 1970 ~ 2002 年包含 4 种资产类的最小方差前沿。那么在子区间 1970 ~ 2002 年的最小方差前沿间会发现什么变动呢？为了了解这一点，我们使用整个区间中的 10 年数据计算样本统计量，而对于每个 10 年求出最小方差前沿。表 11-10 给出了这 4 种资产类在 1970 ~ 1979 年、1980 ~ 1989 年、1990 ~ 2002 年以及整个样本区间的月度资产收益率的样本统计量。

表 11-10 平均收益率、标准差和相关系数矩阵

	标准普尔 500	美国小盘股	MSCI 世界 (除美国外) 指数	美国长期政府债券
A. 平均收益率				
时间段				
1970 ~ 1979 年	7.0%	14.4%	11.8%	5.7%
1980 ~ 1989 年	17.6%	16.7%	21.1%	12.9%
1990 ~ 2002 年	10.4%	13.2%	2.7%	9.9%
总时间段	11.6%	14.6%	11.1%	9.6%
B. 标准差				
时间段				
1970 ~ 1979 年	15.93%	26.56%	16.68%	8.25%
1980 ~ 1989 年	16.41%	19.17%	17.32%	14.19%
1990 ~ 2002 年	15.27%	20.71%	16.93%	8.31%
整个时间段	15.83%	22.18%	17.07%	10.44%
C. 相关系数矩阵				
1970 ~ 1979 年				
标准普尔 500	1			
美国小盘股	0.787	1		
MSCI 世界 (除美国) 指数	0.544	0.490	1	
美国长期政府债券	0.415	0.316	0.218	1
1980 ~ 1989 年				
标准普尔 500	1			
美国小盘股	0.844	1		
MSCI 世界 (除美国) 指数	0.512	0.483	1	
美国长期政府债券	0.310	0.171	0.229	1
1990 ~ 2002 年				
标准普尔 500	1			
美国小盘股	0.611	1		
MSCI 世界 (除美国) 指数	0.647	0.473	1	
美国长期政府债券	0.097	-0.047	0.017	1
整个时间段				
标准普尔 500	1			
美国小盘股	0.731	1		
MSCI 世界 (除美国) 指数	0.573	0.475	1	
美国长期政府债券	0.266	0.138	0.155	1

资料来源：伊博森协会

正如我们所预期的，在子样本区间中所有资产类的样本均值、方差和协方差都有所不同。最初，相关系数在时间上相对稳定。例如，标准普尔 500 和 MSCI 世界非美国指数的 1970 ~ 1979 年、1980 ~ 1989 年和 1990 ~ 2002 年的相关系数分别为 0.544、0.512 和 0.647。与平均收益率的排序不同，由标准差排序的资产类顺序在每个 10 年都是相同的，即美国小盘股是风险最大的资

产类，而债券是风险最小的。我们可以使用统计推断来探求这些跨期数据间差别的原因。然而，在这些初始印象的情况下，让我们首先来考察每个10年的最小方差前沿。

图11-11给出了使用表11-10中1970~1979年，1980~1989年、1990~2002年和整个样本区间的历史收益率统计量所计算出的最小方差前沿。正如这个图中所示，最小方差前沿会在不同区间变化很大。例如，注意到基于1970~1979年的最小方差前沿和基于1980~1989年的最小方差前沿完全不相交。

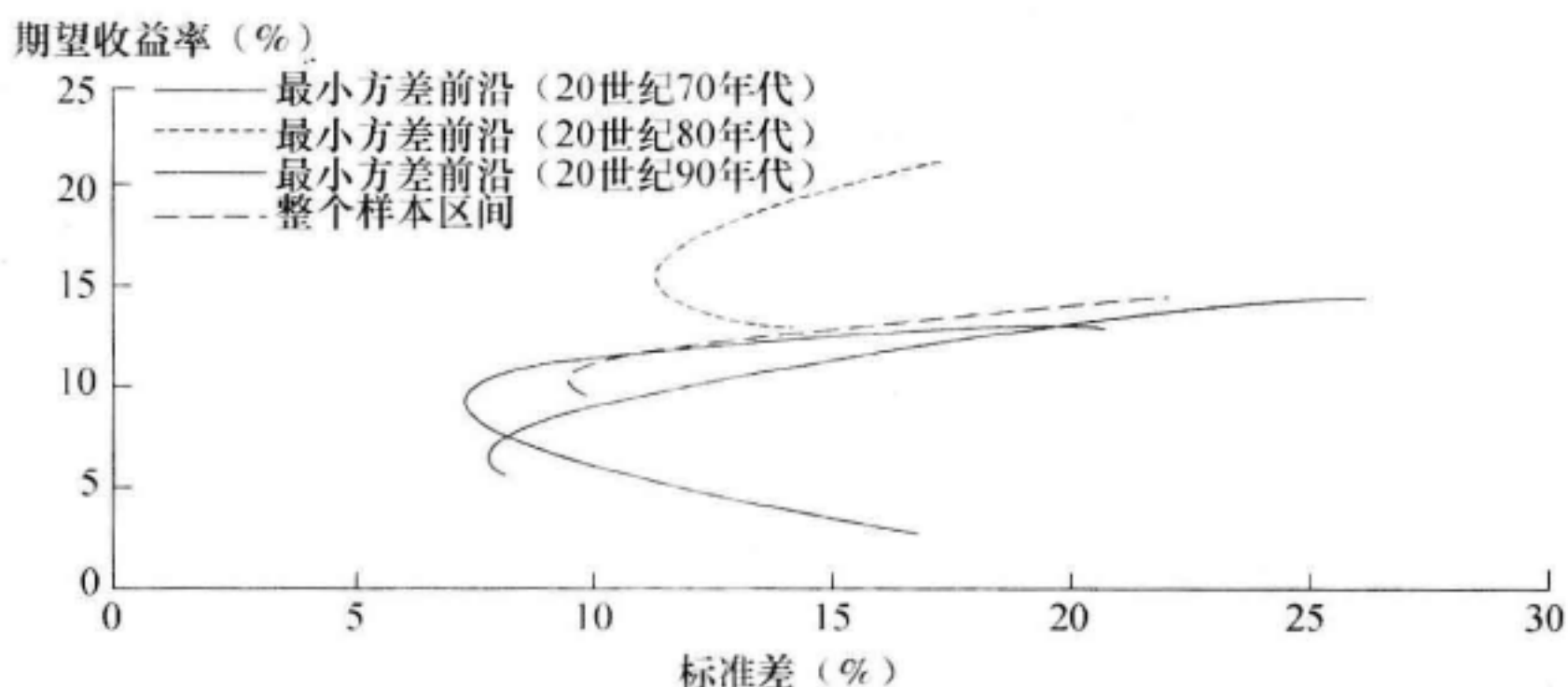


图11-11 历史最小方差前沿的比较

正如之前所提到的，研究者们已经提出了解决投资组合经理所面临的最小方差前沿不稳定性问题的不同方法。

例11-9 耶鲁大学捐赠基金是如何使用均值一方差分析的

戴维·斯文森 (David Swensen) 是耶鲁大学的一位投资主管 (他同样也在耶鲁大学教授投资组合管理)。他在他的书中写到“无约束的均值一方差 (最优化问题) 给出的结果一般不能认为是合理的投资组合……因为该过程涉及重要的简化假设，所以采纳由均值一方差最优化问题所给出的无约束资产配置的点估计没有任何意义。”^①

斯文森的评论强调了从业者所关注的标准均值一方差最优化问题的有用性。在均值一方差分析的假设中，最重要的假设是资产组合中资产的均值、方差和协方差都是已知的。因为最优化过程试图充分利用变量间的微小变化，并且估计的均值和其他参数是存在不确定性的，所以该简化假设会对结果产生重大的影响。正如上面所提到的，解决不稳定性方法包括加入关于资产权重的约束以及修改输入变量的历史样本估计值。尽管斯文森对该方法提出了批评，但是耶鲁大学仍使用均值一方差分析来配置其资产组合；只不过耶鲁大学投资办公室在模型中增加了对于权重的约束，并且没有使用原始的历史输入变量。

11.4 多因素模型

前面我们已经对于市场模型进行了讨论，该模型在历史上是首个用来描述引起资产收益率变动过程的模型。市场模型假设资产收益率的所有可解释的变动都与一个因素 (即市场收益率) 有

^① 参见 Swensen (2000)。

关。但是，资产收益率可能与除了市场收益率之外的其他因素也有关系，例如利率的变动、通货膨胀率或者特定的行业收益率。实际上，投资从业者们目前已经在投资组合管理、风险分析以及投资组合业绩评价中使用多因素模型很多年了。

多因素模型之所以能在投资组合管理业界获得重要的地位，其原因主要有2点：首先，多因素模型能够比市场模型更好地解释资产收益率。[⊖]其次，多因素模型比单因素模型更能提供了一个更加具体的风险分析方法。更加详细的模型无论在被动管理还是积极管理中都是非常有用的。

- **被动管理。**在管理一个旨在跟踪包含许多成分股的指数基金时，投资组合经理可能需要选择指数中的一个证券样本。分析师可以使用多因素模型来匹配指数基金因素的风险暴露与所跟踪指数的风险暴露。
- **积极管理。**在形成刻画证券和投资组合的期望收益率和风险的投资组合时，多因素模型会被使用。许多定量投资管理者会依靠多因素模型来预测阿尔法（超额风险调整后的收益率）或者相对收益率（一个资产或者资产类相对于另一个资产或者资产类的收益率），并使其作为许多不同积极投资策略中的一部分。在评估投资组合时，分析师会使用多因素模型来了解积极投资管理者的收益的来源，以及评估相对于管理者设定的基准（benchmark）（用来比较的投资组合）所承担的风险。

在下面的章节中，我们将解释一些因素模型的基本原理，并讨论一些不同类型的模型以及它们的具体应用。我们也将给出由罗斯（1976）提出的套利定价理论，该理论将投资的期望收益率和它们关于一系列因素的风险相联系。

11.4.1 因素和多因素模型的类型

在本节的开始，我们先定义一些术语，一个因素（factor）是指一个与许多变量相关的共同的或者基本的元素。例如，市场因素就是一个与各个股票收益率相关的基本因素。我们寻找系统因素（systematic factors），该因素影响许多不同资产的平均收益率。这些因素代表定价风险（priced risk），承担该风险的投资者会要求额外的收益。因此，系统因素应该能帮助解释收益率。

许多不同的多因素模型已经被人们所提出与研究。我们可以根据所使用的因素类型将它们的绝大部分归为3个主要类别：

- 在宏观经济因素模型（macroeconomic factor models）中，因素是宏观经济变量中能够显著解释股票收益率的意外值。因素可以被认为是影响公司未来期望现金流或者将这些现金流折现到当前时刻所使用的利率。
- 在基本面因素模型（fundamental factor models）中，因素是能解释股价横截面差异的股票或者公司的重要特征。常用的基本面因素有账面价值—股价比、市值、市盈率和财务杠杆。
- 在统计因素模型（statistical factor models）中，统计方法被应用于一组历史收益率以确定出可以用下面2种方法中的一种解释其历史收益率的投资组合。在因素分析模型中，因素是可以最佳解释（复制）历史收益率协方差的投资组合。而在主成份模型中，因素是

⊖ 例如，参见 Burmeister 和 McElroy（1988）。这2位作者给出了在1%显著性水平下，CAPM可以被拒绝，而支持一个含有多个因素的无套利定价利率模型。我们将在本章的后面讨论无套利定价理论。

可以最佳解释历史收益率方差的投资组合。

一些实际使用的因素模型可能会包含上述分类中的不止一种特征。我们将这类模型称为混合因素模型 (mixed factor models)。

我们的讨论将主要关注于宏观经济因素模型和基本面因素模型。从业界的使用程度上来看,他们一般也更喜欢基本面因素模型和宏观经济因素模型,这可能是因为这类模型更加容易解释;但是,统计因素模型也有其支持者,并且也被应用于实际投资之中。

11.4.2 宏观经济因素模型的结构

给出资产收益率表达式的宏观经济因素模型假设:各个资产的收益率只与总体经济相关的一些因素有关,例如通货膨胀或者实际产出。^①我们一般可以将意外值 (surprise) 定义为实际值减去预测 (或者期望) 值。一个因素的意外值是该因素无法预测收益率的部分,而因素的意外值形成了模型中的自变量。该想法与基本面因素模型中表示为自变量的收益率不同 (与收益率的意外值不同),或者与市场模型中的收益率也不同。

假如有 K 个可以解释资产收益率的因素。那么在一个宏观经济因素模型中,可以用下面的等式来表示资产 i 的收益率:

$$R_i = a_i + b_{i1}F_1 + b_{i2}F_2 + \cdots + b_{iK}F_K + \epsilon_i \quad (11-16)$$

式中, R_i 表示资产 i 的收益率; a_i 表示资产 i 的期望收益率; F_k 表示因素 k 的意外值, $k=1, 2, \cdots, K$; b_{ik} 表示资产 i 的收益率对于因素 k 意外值的敏感性, $k=1, 2, \cdots, K$; ϵ_i 表示资产 i 收益率不能被因素模型所解释的部分,其均值为0。

我们所说的宏观经济因素中的意外值到底意味着什么呢?假如我们分析股票的月度收益率。在每个月月初,我们已经预测出了该月的通货膨胀率。该预测可能来自于一个经济计量模型或者一个职业经济预测家。假如我们在月初的预测为该月的通货膨胀率为0.4%。在该月月底,我们发现该月的实际通货膨胀率为0.5%。在任何一个月,

$$\text{实际通货膨胀率} = \text{预测通货膨胀率} + \text{通货膨胀率意外值}$$

在这个例子中,实际通货膨胀率为0.5%,而预测的通货膨胀率为0.4%。因此,通货膨胀率的意外值为 $0.5 - 0.4 = 0.1\%$ 。

以意外值定义的因素所造成的影响是什么呢?假如我们认为通货膨胀率和国内生产总值 (GDP) 增长率是定价风险。(GDP 是一个以货币度量的反映一个国家境内所生产的商品与服务的指标。) 我们没有使用这些变量的预测值,这是因为预测值已经反映在股票价格以及它们的期望收益率之中。截距 a_i , 资产 i 的期望收益率,反映宏观经济变量的预测值对于期望股票收益率的影响。另一方面,该月宏观经济变量中的意外值包含了关于该变量新的信息。因此,该模型的设定将一个资产的收益率解释为3个部分:资产的期望收益率,由于关于因素的新信息而不可预期的收益率以及一个误差项。

考虑一个因素模型,其中每种资产的收益率都与2个因素相关。例如,我们可能假设一个特定股票的收益率与利率的意外值和 GDP 增长率的意外值有关。对于股票 i , 该股票的收益率可以被刻画为

$$R_i = a_i + b_{i1}F_{INT} + b_{i2}F_{GDP} + \epsilon_i \quad (11-17)$$

^① 例如,参见 Burmeister, Roll 和 Ross (1994)。

式中, R_i 表示股票 i 的收益率; a_i 表示股票 i 的期望收益率; b_{i1} 表示股票 i 的收益率对于利率意外值的敏感性; F_{INT} 表示利率意外值; b_{i2} 表示股票 i 的收益率对于 GDP 增长率意外值的敏感性; F_{GDP} 表示 GDP 增长率意外值; ϵ_i 表示资产 i 收益率不能被因素模型所解释的部分, 其均值为 0。

首先我们考虑如何解释 b_{i1} 。因素模型预测每一个百分点的利率意外值将带来 b_{i1} 百分比的股票 i 收益率。斜率系数 b_{i2} 具有与 GDP 增长率因素的类似解释。因此, 斜率系数很自然地解释为对资产因素的敏感性。[⊖] 对一个因素敏感性 (factor sensitivity) 是一个度量在保持其他因素不变的情况下, 每单位因素的增长所造成的收益率的增长。

现在考虑如何来解释截距 a_i 。回忆一下, 误差项的均值或者平均值为 0。如果利率和 GDP 增长率的意外值都是 0, 因素模型预测资产 i 的 i 的收益率将为 a_i 。因此, a_i 是股票 i 的期望收益率。

最后, 考察误差项 ϵ_i 。截距 a_i 表示资产的期望收益率。值 $(b_{i1}F_{INT} + b_{i2}F_{GDP})$ 表示由因素意外值带来的收益率, 而我们已经将这些解释为其他资产所共同的风险来源。项 ϵ_i 是收益率中不能由期望收益率或者因素意外值所解释的部分。如果我们有充分反映了共同风险 (因素) 的来源, 那么 ϵ_i 一定代表特定资产的风险。对于一只股票, 它可能代表由非预期的特定公司事件所带来的收益率。

当我们介绍套利定价理论时, 我们将进一步讨论期望收益率。在宏观经济因素模型中, 因素意外值的时间序列是最先获得的。我们使用回归分析来估计资产对于因素的敏感性。在我们的讨论中, 我们假设你没有亲自估计敏感性和截距; 相反, 你可以使用从其他来源 (例如, 专门使用因素模型的许多咨询公司之一) 获得的估计值。[⊗] 当我们有了投资组合中每种资产参数之后, 我们可以通过每种资产参数的加权平均计算出投资组合的参数。在计算中每种资产的权重等于投资组合市场总值中每种资产所占的比例。

例 11-10 2 只股票投资组合的因素敏感性

假如股票收益率受到 2 个共同因素的影响: 通货膨胀率的意外值和 GDP 增长率的意外值。一位投资经理正在分析包含 2 只股票 Manumatic (MANM) 和 Nextech (NXT) 的投资组合的收益率。下面的等式描述了这 2 只股票的收益率, 其中因素 F_{INFL} 和 F_{GDP} 分别表示通货膨胀率和 GDP 增长率的意外值:

$$R_{MANM} = 0.09 - 1F_{INFL} + 1F_{GDP} + \epsilon_{MANM}$$

$$R_{NXT} = 0.12 + 2F_{INFL} + 4F_{GDP} + \epsilon_{NXT}$$

在评价这 2 个等式中通货膨胀率和 GDP 增长率的意外值时, 以百分比表示的数量需要转化为以小数表示的数量。1/3 的投资组合被投资于 Manumatic 股票, 而 2/3 被投资于 Nextech 股票。

⊖ 因素的敏感性有时也被称为因素贝塔 (factor betas) 或者因子载荷 (factor loadings)。

⊗ 如果你想要亲自估计宏观经济因子模型, 那么就要采用下列步骤。首先, 估计每个宏观经济意外值的时间序列 (例如, 你可能使用每个不同宏观经济序列的时间序列模型的残差)。接着, 使用时间序列数据来将特定资产的收益率关于不同宏观经济因素的意外值进行回归。

(1) 给出投资组合收益率的表达式。

(2) 求出投资组合的期望收益率。

(3) 给定通货膨胀率的意外值和 GDP 增长率的意外值分别为 1% 和 0%，计算出投资组合的收益率。(假定 MANM 和 NXT 的误差项都等于 0.5%。)

解：

(1) 投资组合的收益率为如下 2 只股票收益率的加权平均数：

$$\begin{aligned} R_p &= (1/3)(0.09) + (2/3)(0.12) + [(1/3)(-1) + (2/3)(2)]F_{\text{INFL}} \\ &\quad + [(1/3)(1) + (2/3)(4)]F_{\text{GDP}} + (1/3)\epsilon_{\text{MANM}} + (2/3)\epsilon_{\text{NXT}} \\ &= 0.11 + 1F_{\text{INFL}} + 3F_{\text{GDP}} + (1/3)\epsilon_{\text{MANM}} + (2/3)\epsilon_{\text{NXT}} \end{aligned}$$

(2) 投资组合的期望收益率为 11%，该值可以由第 1 问中表达式的截距值得到。

$$\begin{aligned} (3) \quad R_p &= 0.11 + 1F_{\text{INFL}} + 3F_{\text{GDP}} + (1/3)\epsilon_{\text{MANM}} + (2/3)\epsilon_{\text{NXT}} \\ &= 0.11 + 1(0.01) + 3(0) + (1/3)(0.05) + (2/3)(0.005) \\ &= 0.125 \text{ 或者 } 12.5\%。 \end{aligned}$$

11.4.3 套利定价理论和因素模型

在 20 世纪 70 年代，史蒂芬·罗斯创建了套利定价理论 (arbitrage pricing theory, APT)，并将其作为 CAPM 的一个替代理论。APT 将一个资产（或者投资组合）的期望收益率描述为一个由一组因素表示的资产（或者投资组合）风险的线性函数。类似于 CAPM，APT 描述了一个金融市场均衡。然而，APT 所做的假设要弱于 CAPM。APT 依赖于这样 3 个假设：

1. 一个因素模型能够用来描述所有资产收益率。
2. 存在许多资产，从而投资者可以获得消除特定资产风险的充分分散化的投资组合。
3. 在充分分散化的投资组合中不存在套利机会。

套利 (arbitrage) 是指一种无需任何净投资资金而获得期望正的净收益率的无风险的交易策略。^①一个套利机会 (arbitrage opportunity) 是指执行套利的机会——一个无需任何净投资资金而获得期望正的净收益率的机会。

在第 1 个假设中，因素的数量没有明确给出。第 2 个假设允许我们形成具有因素风险的投资组合，但允许存在特定资产风险。第 3 个假设是金融市场均衡的条件。

实证的证据表明假设 2 是合理的。当一个投资组合包含许多股票时，特定资产或者单只股票的非系统性风险几乎对于投资组合收益率的方差没有任何影响。罗尔和罗斯 (2001) 发现只有充分分散化的投资组合的方差中只有 1% ~ 3% 来自于投资组合中个股的非系统性方差，正如下面的图 11-12 所示。

① 正如我们将看到的，套利通常涉及使用卖空另一项资产所获得的资金来投资于某项资产，因此，从净值上来看，没有资金被投入。卖空就是指卖出借来的资产。注意“套利”这个词有时也用来描述在显著风险存在情况下的投资操作。

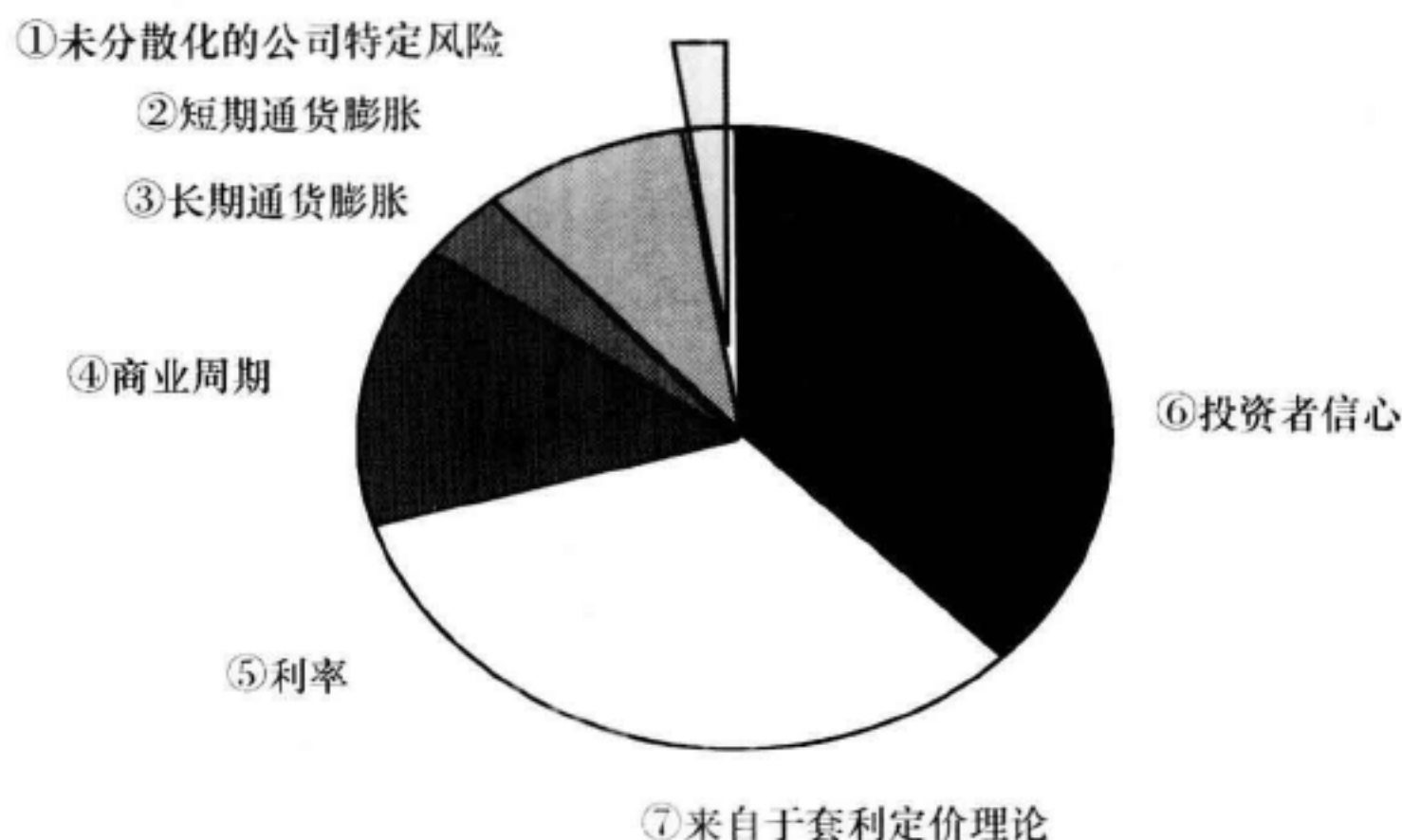


图 11-12 波动率的来源：一个充分分散化投资组合的例子

资料来源：2001 年 5 月 25 日，万维网：www.rollross.com/apr.html。

根据 APT，如果上面 3 个假设成立，那么下列等式成立：^①

$$E(R_p) = R_F + \lambda_1 \beta_{p,1} + \cdots + \lambda_K \beta_{p,K} \quad (11-18)$$

式中， $E(R_p)$ 表示投资组合 p 的期望收益率； R_F 表示无风险利率； λ_j 表示因素 j 的风险溢价； $\beta_{p,j}$ 表示投资组合对于因素 j 的敏感性； K 表示因素的数目。

APT 等式，即式 (11-18)，认为任何充分分散化的投资组合的期望收益率都与其因素敏感性线性相关。^②

因素风险溢价 (factor risk premium) (或者因素价格 (factor price)) λ_j 表示一个对于因素 j 具有 1 的敏感性，而其他因素的敏感性都等于 0 的投资组合超出无风险收益率的期望收益率。该投资组合也被称为因素 j 的纯因素投资组合 (pure factor portfolio)。

例如，假设我们具有一个对于因素 1 具有敏感性为 1 而对于其他因素的敏感性为 0 的投资组合。在 E_1 表示该投资组合期望收益率的情况下，式 (11-18) 表明该投资组合的期望收益率为 $E_1 = R_F + \lambda_1 1$ ，因此， $\lambda_1 = E_1 - R_F$ 。假如 $E_1 = 0.12$ 而 $R_F = 0.04$ 。那么因素 1 的风险溢价为 $\lambda_1 = 0.12 - 0.04 = 0.08$ 或者 8%。对于第 1 个因素敏感性 1 的增长我们将获得期望收益率 8 个百分点的增长。

APT 等式和多因素模型等式 (式 (11-17)) 之间的关系是什么呢？在讨论多因素模型中，我们认为截距项是投资的期望收益率。而 APT 等式解释该投资的期望收益率是在均衡的时候达到的。因此，如果 APT 成立，它将对于多因素模型中的截距项施加限制，因为 APT 模型告诉我们截距项的值应该是什么。例如，在例 11-10 中，APT 将模型中 MANA 收益率的截距 0.09 解释为在给定股票对于通货膨胀率和 GDP 增长率因素敏感性和这些因素风险溢价条件下的期望收益率。我们事实上能够将 APT 等式代入多因素模型来获得被称为以收益率形式表示的 APT 模型。^③

为了使用 APT 等式，我们需要估计其参数。APT 等式的参数是无风险收益率和因素风险溢价

① 假设存在一项无风险资产。如果不存在无风险资产，那么我们将用 λ_0 代替 R_F 来表示一项对于所有因素具有 0 敏感性的风险投资组合。虽然因素的数量没有明确给出，但是它必须远低于资产的数量，这是在实际中需要满足的一个条件。

② APT 等式也能将投资的期望收益率描述为 (至少近似描述为) 在确定性条件下特定资产的风险。

③ 一个有趣的问题是 APT 和 CAPM 之间的关系。如果市场是单因素模型中的因素，那么 APT (式 (11-18)) 就和 CAPM 一致。如果 APT 模型中的风险溢价满足一定的限制，CAPM 也可以与多个因素的 APT 模型一致；这些与 CAPM 有关的限制条件已经在统计检验中被多次拒绝了。例如，参见 Burmeister 和 McElroy (1988)。

(因素敏感性是针对各项投资的)。例 11-11 给出了单因素假设下一组投资组合的期望收益率和因素敏感性是如何来确定 APT 模型中的参数的。

例 11-11 确定单因素 APT 模型中的参数

假设我们有 3 个充分分散化的投资组合，其中每个都对于相同的一个因素敏感。表 11-11 给出了这些投资组合的期望收益率和因素的敏感性。假设期望收益率是反映了一年投资区间的情况。

表 11-11 单因素模型的样本投资组合		
投资组合	期望收益率	因素敏感性
A	0.075	0.5
B	0.150	2.0
C	0.070	0.4

我们可以使用这些数据来确定 APT 等式中的参数。根据式 (11-18)，对于任一充分分散化的投资组合，并且假设可以解释收益率的只有一个因素，我们有 $E(R_p) = R_f + \lambda_1 \beta_{p,1}$ 。因素敏感性和期望收益率是已知的；因此，存在 2 个未知变量，参数 R_f 和 λ_1 。因为两点确定一条直线，所以我们需要建立 2 个等式。选择投资组合 A 和 B，我们有

$$E(R_A) = 0.075 = R_f + 0.5\lambda_1$$

以及

$$E(R_B) = 0.150 = R_f + 2\lambda_1$$

从投资组合 A 的等式中，我们有 $R_f = 0.075 - 0.5\lambda_1$ 。将这个无风险收益率的表达式代入投资组合 B 的等式中，有

$$0.15 = 0.075 - 0.5\lambda_1 + 2\lambda_1$$

$$0.15 = 0.075 + 1.5\lambda_1$$

因此，我们有 $\lambda_1 = (0.15 - 0.075)/1.5 = 0.05$ 。将 λ_1 的值带回到投资组合 A 期望收益率的等式中，得到

$$0.075 = R_f + 0.05 \times 0.5$$

$$R_f = 0.05$$

因此，无风险收益率为 0.05 或者 5%，共同因素的因素溢价也为 0.05 或者 5%。APT 等式为

$$E(R_p) = 0.05 + 0.05\beta_{p,1}$$

投资组合 C 具有 0.4 的因素敏感性。相应地，如果不存在套利机会， $0.05 + (0.05 \times 0.4) = 0.07$ 或者 7%。有表 11-11 所给出的投资组合 C 的期望收益率为 7%。因此，在这个例子中不存在套利机会。

例 11-12 检验投资组合收益率与无套利是否相一致

在这个例子中，我们给出如何通过检验是否存在套利机会来说明是否充分分散化投资组合的一组期望收益率是与 APT 相一致的。在例 11-11 中，我们得到 3 个投资组合的期望收益率和因素敏感性是与单因素 APT 模型 $E(R_p) = 0.05 + 0.05\beta_{p,1}$ 相一致的。假如我们扩大投资组合集，使其包含第 4 个充分分散化的投资组合，投资组合 D。表 11-12 重复给出了表 11-11 中投资组合 A、B 和 C 的数据，另外还提供了投资组合 D 的数据，以及一个由投资组合 A 和 C 形成的投资组合。

表 11-12 单因素模型的样本投资组合

投资组合	期望收益率	因素敏感性	投资组合	期望收益率	因素敏感性
A	0.075 0	0.50	D	0.080 0	0.45
B	0.150 0	2.00	0.5 A + 0.5 C	0.072 5	0.45
C	0.070 0	0.40			

投资组合的期望收益率和因素敏感性是投资组合中各个资产期望收益率和因素敏感性的加权平均值。假设我们构建一个包含 50% 投资组合 A 和 50% 投资组合 C 的投资组合。表 11-12 给出了该投资组合的期望收益率等于 $(0.5)(0.0750) + (0.5)(0.07) = 0.0725$ 或者 7.25%。该投资组合的因素敏感性为 $(0.50)(0.50) + (0.50)(0.40) = 0.45$ 。

套利定价理论假设充分分散化的投资组合不能提供任何套利机会。如果初始投资为 0，而我们没有承担任何风险，那么期末期望的现金流也应该是 0。在这个例子中，关于因素风险的期望收益率的配置提供了一个涉及资产组合 A、C 和 D 的套利机会。投资组合 D 在给定其因素敏感性条件下，提供了一个过高的期望收益率。根据例 11-11 中所估计的 APT 模型，除非 $E(R_D) = 0.05 + 0.05\beta_{D,1} = 0.05 + (0.050 \times 45) = 0.0725$ ，否则就存在套利机会。因此，D 的期望收益率应该为 7.25%。而事实上，D 的期望收益率为 8%。投资组合 D 相对于其因素风险被低估了。我们将买入 D（在投资组合中持有多头头寸）从而利用该套利机会（套利投资组合（arbitrage portfolio））。我们使用卖空由 A 和 C 组成的与 D 具有相同因素敏感度 0.45 的投资组合来购买 D。正如我们如上所示，一个等权重的 A 和 C 的投资组合正好具有 0.45 的因素敏感性。

因此，套利涉及如下策略：投资 10 000 美元到投资组合 D，并通过卖空等权重的投资组合 A 和 C 的投资组合来融得相应资金；然后在一年年末平仓（该投资区间是与期望收益率的时间段一致的）。表 11-13 给出了该套利策略的套利利润。表的最后一行给出了该套利组合的净现金流。

表 11-13 样本投资组合中的套利机会

	初始现金流	最终现金流	因素敏感性
投资组合 D	- \$10 000.00	\$10 800.00	0.45
投资组合 A 和 C	\$10 000.00	- \$10 725.00	-0.45
总和	\$0.00	- \$75.00	0.00

正如表 11-13 所示，如果我们用 10 000 美元购买投资组合 D，并卖出 10 000 美元等权重的投资组合 A 和 C 组成的投资组合，那么我们的初始净现金流为 0 美元。在一年年末，我们投资于投资组合 D 的期望价值为 $\$10\,000(1 + 0.08) = \$10\,800$ 。而一年年末卖空投资组合 A 和 C 的期望价值为 $-\$10\,000(1.0725) = -\$10\,725$ 。所以我们一年持有投资头寸的总期望现金流为 \$75。

那么，风险情况是怎么样的呢？表 11-13 表明因素风险已经被消除了：购买 D 而卖空等权重的投资组合 A 和 C 产生了一个因素敏感性为 $0.45 - 0.45 = 0$ 的投资组合的。该投资组合是充分分散化的，而我们假设任何特定资产风险都是可以忽略的。

因为套利是可能的，所以投资组合 A、C 和 D 不能同时达到相同的均衡。APT 模型的唯一参数集不能同时描述这 3 个投资组合的收益率。如果投资组合 D 实际上的期望收益率为 8%，那么投资者将会抬高其价格直至其期望收益率下降并使得套利机会消除。因此，套利能够回复期望收益率间的均衡关系。

在例 11-11 中，我们说明了单因素 APT 模型参数是如何由数据所确定的。例 11-13 表明如何确定包含多于一个因素的模型中的参数。

例 11-13 确定 2 个因素模型中的参数

假设 2 个因素，通货膨胀率的意外值（因素 1）以及 GDP 增长率的意外值（因素 2），能够解释收益率。根据 APT，除非 $E(R_p) = R_f + \lambda_1 \beta_{p,1} + \lambda_2 \beta_{p,2}$ ，否则就存在一个套利机会。

我们的目标是要估计模型中的 3 个参数 R_f ， λ_1 和 λ_2 。我们也有表 11-14 中 3 个充分分散化投资组合 J、F 和 L 的假设数据。

表 11-14 2 个因素模型的样本投资组合

投资组合	期望收益率	通货膨胀因素的敏感性	GDP 因素的敏感性
J	0.14	1.0	1.5
K	0.12	0.5	1.0
L	0.11	1.3	1.1

如果市场处于均衡状态（不存在套利机会），那么 3 个投资组合的期望收益率应该能被具有相同参数集的 2 个因素 APT 进行描述。使用表 11-14 中的期望收益率和收益率的敏感性可以得到

$$E(R_J) = 0.14 = R_f + 1.0\lambda_1 + 1.5\lambda_2$$

$$E(R_K) = 0.12 = R_f + 0.5\lambda_1 + 1.0\lambda_2$$

$$E(R_L) = 0.11 = R_f + 1.3\lambda_1 + 1.1\lambda_2$$

我们在 3 个等式中具有 3 个未知变量，所以我们可以使用替代法来求解参数。我们首先想要获得具有 2 个未知参数的 2 个等式。求解以 $E(R_J)$ 表示的无风险收益率的等式，

$$R_f = 0.14 - 1.0\lambda_1 - 1.5\lambda_2$$

将该无风险收益率的表达式代入 $E(R_K)$ 的等式，在化简之后，我们求得 $\lambda_1 = 0.04 - \lambda_2$ 。使用 $\lambda_1 = 0.04 - \lambda_2$ 来消去 $E(R_J)$ 等式中的 λ_1 ，

$$0.10 = R_f + 0.5\lambda_2$$

使用 $\lambda_1 = 0.04 - \lambda_2$ 消去 $E(R_L)$ 等式中的 λ_1 ，

$$0.058 = R_f - 0.2\lambda_2$$

马上使用上述 2 个关于 R_f 和 λ_2 的等式，我们就能求得 $\lambda_2 = 0.06$ （我们在这 2 个等式的第 1 个中求出无风险收益率在使用第 2 个等式中的表达式）。因为 $\lambda_1 = 0.04 - \lambda_2$ ， $\lambda_1 = -0.02$ 。最后， $R_f = 0.14 - 1.0(-0.02) - 1.5(0.06) = 0.07$ 。总结得：

$$R_f = 0.07 \text{（风险收益率为 7\%。）}$$

$$\lambda_1 = -0.02 \text{（通货膨胀风险溢价为 -2\% 每单位敏感度。）}$$

$$\lambda_2 = 0.06 \text{（GDP 风险溢价为 6\% 每单位敏感度。）}$$

因此，对于这 3 个投资组合的 APT 等式为

$$E(R_p) = 0.07 - 0.02\beta_{p,1} + 0.06\beta_{p,2}$$

这个例子说明了确定 APT 模型参数的计算方法。它也表明了一个因素的风险溢价实际可能为负。

在例 11-13 中, 我们计算出的通货膨胀因素的风险溢价为负的。对于该负的通货膨胀风险溢价的一种解释为大部分股票具有对于通货膨胀风险负的敏感性 (它们的收益率倾向于在一个正的通货膨胀意外值出现时下降)。一个具有正的通货膨胀敏感性的资产将作为一个对冲通货膨胀的资产; 与通货膨胀风险的因素投资组合相关的溢价结果可能为负。

11.4.4 基本面因素模型的结构

我首先给出宏观经济因素模型

$$R_i = a_i + b_{i1}F_1 + b_{i2}F_2 + \cdots + b_{iK}F_K + \epsilon_i$$

我们也能用该等式来表示基本面因素模型, 但是我们需要对于模型中的各项进行不同的解释。

在基本面因素模型中, 因素被称为收益率而不是收益率关于预期值的意外值, 所以它们的期望值一般不为 0。该方法改变了我们对于截距的解释, 我们不再将其解释为期望收益率。^①

在大部分基本面因素模型中, 我们对于因素敏感性的解释也将不同。在基本面因素模型中, 因素敏感性是证券的属性。考虑股票的一个具有红利收益率因素的基本面模型。一个资产相对于红利因素的敏感性是该属性本身的值, 即其红利收益率; 敏感性通常是经过标准化的。具体地说, 一项资产 i 相对于给定因素的敏感性将计算为资产属性的值减去所有股票属性的平均值, 再除以所有股票属性的标准差。标准化贝塔 (Standardized beta) 为

$$b_{ij} = \frac{\text{资产 } i \text{ 属性的值} - \text{平均属性的值}}{\sigma(\text{属性的值})} \quad (11-19)$$

继续我们关于红利收益率的例子, 在标准化之后, 一个具有平均红利收益率的股票所具有的因素敏感性为 0, 具有高于平均红利收益率一个标准差的股票所具有的因素敏感性为 1, 而具有低于平均红利收益率一个标准差的股票所具有的因素敏感性为 -1。例如, 假设一项投资的红利收益率为 3.5%, 而所有考虑的股票的平均红利收益率为 2.5%。并且进一步假设所有股票红利收益率的标准差为 2%。那么该投资相对于红利收益率的敏感性为 $(3.5\% - 2.5\%) / 2\% = 0.50$, 或者是高于平均值半个标准差。尽管变量度量的方式和单位都是不同的, 但是该度量标准允许对所有因素敏感性以相同的方式进行解释。该解释的一个例外情况是对于二元取值变量的因素, 如是否属于某个行业。一家公司要么属于这个行业, 要么不属于这个行业。行业因素的敏感性将是 0-1 的虚拟变量; 在刻画频繁从事于多个行业公司的模型中, 敏感性的值将对于该公司所从事的每个行业都取值为 1。^②

宏观经济多因素模型与基本面因素模型的第 2 个区别是: 在前者中, 我们先得到因素 (意外值) 序列, 然后再通过回归估计出因素的敏感性; 而对于后者, 我们一般先设定因素的敏感性 (属性) 再通过回归估计出因素的收益率。^③

11.4.5 多因素模型在当前实践中的运用

在前一节中, 我们解释了多因素模型和 APT 的一些基本概念。现在我们开始讨论在实际业界

① 如果系数没有经过之后段落所述的标准化, 那么截距就可能被解释为无风险收益率, 因为它将是没有因素风险 (零因素贝塔) 和特定资产风险的资产的收益率。在标准化系数之后, 截距仅仅解释为包含在回归中的一个截距项, 因此, 期望特定资产的风险等于 0。

② 为了进一步解释 0-1 变量, 是否属于某个行业是以一个名义度量指标来刻画的, 因为我们可以说出一家公司属于哪个行业。一个名义变量可以在回归中表示为一个虚拟变量 (一个只取 0 或 1 的变量)。更多关于虚拟变量的内容, 参见多元回归一章。

③ 在一些归类为基本面的模型中, 因素敏感性指的是回归系数, 并且它们不是事先定义的。

所使用的一些模型。

11.4.5.1 宏观经济因素模型

陈 (Chen)、罗尔 (Roll) 和罗斯 (Ross) (1986) 首先提出了宏观经济因素模型。从统计上的研究表明多个因素在对于美国股票的平均收益率的解释时非常重要, 根据这一结果, 陈等人建议使用相对较少的一组宏观因素来作为影响美国股票市场的主要因素。陈等人在研究中所采用的因素为: (1) 通货膨胀率, 包括非预期的通货膨胀率和期望通货膨胀率的变动; (2) 与利率期限结构有关的一个因素, 用长期政府债券收益率减去一个月国库券收益率来表示; (3) 反映市场风险和投资者风险厌恶的一个因素, 用低评级债券和高评级债券的收益率之差来表示; (4) 工业产值的变动。

任一用来解释资产收益率的因素的有用性一般都是经过历史数据进行过检验的。如果我们能够对于一个因素为什么在解释平均收益率中非常重要的这个问题给出一个经济的解释, 那么我们就能够相信一个因素将能解释未来的收益率增长。我们有充分的理由来解释陈 (Chen) 等人提出的所有这4个因素。例如, 通货膨胀不仅会影响经营的现金流, 还会影响投资者用来折现这些现金流的贴现率。工业产值的变动会影响经营的现金流和投资者所面对的投资机会。例 11-14 具体给出了当前在陈 (Chen) 等人模型的基础上所使用的宏观经济因素模型。

例 11-14 一个宏观经济因素模型中的期望收益率

布曼斯特 (Burmeister)、罗尔 (Roll) 和罗斯 (Ross) (1994) 提出了一个解释美国股票收益率的宏观经济因素模型。该模型被简称为 BIRR 模型。BIRR 模型包括下面 5 个因素:

1. 信用风险: 有风险的公司债券与政府债券收益率之差的非预期变动 (2 个债券的到期期限都为 20 年)。有风险的公司债券比政府债券拥有更大的违约风险。投资者对于该风险的态度应该会影响股票的平均收益率。该因素名称的含义是当投资者对债券的信用很有信心时, 投资者会愿意以更小的收益来承担这种风险。
2. 投资期限风险: 20 年政府债券和 30 天国库券收益率之差的非预期变化。该因素反映投资组合愿意投资于长期债券的程度。
3. 通货膨胀风险: 通货膨胀率的非预期变化。几乎所有的股票具有对于该因素负的风险暴露, 即随着通货膨胀的正意外值, 它们的收益率会下降。
4. 商业周期风险: 实际商业活动水平的非预期变化。一个正的意外值或者非预期的变动表明经济的预期收益率 (以常数美元度量的) 已经增加。
5. 市场择机风险: 标准普尔 500 总收益率不能由前 4 个风险因素所解释的部分。[⊖] 几乎所有的股票具有对于该因素的正的敏感性。

前 4 个因素与陈等人根据他们试图刻画的经济影响所给出的因素非常类似。第 5 个因素认为资产定价中存在着围绕着给定变量正确值的不确定性; 该因素刻画了标准普尔 500 收益率不能由前 4 个因素所解释的影响。

标准普尔 500 是一个广泛使用的包含美国主要行业中 500 只主要公司股票的指数。布曼斯特 (Burmeister) 等人使用标准普尔 500 指数来度量 5 个因素对于标准普尔 500 平均超额收益率 (即超过国库券收益率的部分) 的影响。表 11-15 给出了他们的结果。

⊖ 由于因素构建的方式, 标准普尔 500 本身关于市场择机风险的敏感性为 1。

表 11-15 对于标准普尔 500 年度期望超额收益率的解释

风险因素	因素敏感性	风险溢价 (%)	因素对于期望收益率的影响 (%)
信用风险	0.27	2.59	0.70
投资期限风险	0.56	-0.66	-0.37
通货膨胀风险	-0.37	-4.32	1.60
商业周期风险	1.71	1.49	2.55
市场择机风险	1.00	3.61	3.61
期望超额收益率			8.09

资料来源：布曼斯特等人

估计出的 APT 模型为 $E(R_p) = \text{国库券收益率} + 2.59 \text{ (信用风险)} - 0.66 \text{ (投资期限风险)} - 4.32 \text{ (通货膨胀风险)} + 1.49 \text{ (商业周期风险)} + 3.61 \text{ (市场择机风险)}$ 。该表表明标准普尔 500 对于除通货膨胀风险以外的每个风险因素都具有正的风险暴露。对于超额收益率影响最大的 2 个因素是市场择机风险和商业周期风险。根据该表的数据，这个模型预测标准普尔 500 将具有 8.09% 的超过国库券收益率的期望超额收益率。因此，如果 30 天的国库券收益率为 4%，那么预测的标准普尔 500 收益率将为一年 $4 + 8.09 = 12.09\%$ 。

在例 11-15 中，我们将举例说明如何使用布曼斯特 (Burmeister) 等人的因素模型来估计一位管理美国积极核心股票投资组合（一个积极管理的投资于大盘股的投资组合）的投资组合经理所确定的因素值。

例 11-15 总体经济的风险暴露

威廉姆斯·惠更斯 (William Hughes) 是负责一个美国核心股票投资组合的一位投资组合经理，该投资组合评价的比较基准为标准普尔 500 指数。因为惠更斯的业绩表现将与该基准进行比较，所以了解惠更斯相对于标准普尔 500 指数所承担的积极因素是非常有用的。在关注总体经济风险暴露的同时，我们使用布曼斯特 (Burmeister) 等人已经提出的模型。表 11-16 给出了惠更斯的数据。

我们看到投资组合经理完全根据信用风险和投资期限风险来跟踪标准普尔指数收益率，但是相对来说，更加偏向于商业周期风险。另外，该投资组合具有较小的相对于市场择机因素的正的超额风险暴露。

我们可以使用相对于商业周期风险的超额风险暴露来解释超额敏感性的数值。忽略非系统性风险并保持其他因素不变，如果商业周期因素有一个 1% 的意外值，那么我们预计投资组合收益率将会比标准普尔 500 收益率高 0.0054 或者 0.54%。相反，对于商业周期 -1% 的冲击，我们预计投资组合收益率将会比标准普尔 500 收益率低 0.0054 或者 0.54%。

因为超额风险暴露为 0.54，投资组合经理似乎会猜测经济处于扩张期（相对于基准）。如果该因素的猜测是错误的，那么惠更斯将可能承担不想要的风险。如果他做出了该猜测，那么该猜测的理由是什么呢？

表 11-16 一个核心股票投资组合的超额因素敏感性

风险因素	核心投资组合的因素敏感性	标准普尔 500 的因素敏感性	核心投资组合的超额因素敏感性
信用风险	0.27	0.27	0.00
投资期限风险	0.56	0.56	0.00
通货膨胀风险	-0.12	-0.37	0.25
商业周期风险	2.25	1.71	0.54
市场择机风险	1.10	1.00	0.10

我们在解释投资组合经理相对于通货膨胀因素的超额敏感性 0.25 时必须要小心。标准普尔 500 具有负的通货膨胀因素的风险暴露。0.25 这个值代表核心投资组合更小的负的通货膨胀风险暴露,即更小而不是更多的通货膨胀风险。从表 11-16 中我们可以看到因为通货膨胀风险的风险溢价为负,所以惠更斯通过对于通货膨胀率的猜测而放弃相对于基准的期望收益率。那么,他对于通货膨胀因素猜测的依据是什么呢?

市场择机因素具有与 CAPM 关于一只股票倾向于对于总体市场产生反映的贝塔类似的解释,即在所有其他条件不变的情况下,值越高反映市场收益率具有更高的敏感性。但是市场择机因素只反映标准普尔 500 指数收益率中不能由其他 4 个因素解释的部分,所以这 2 个概念是不同的。虽然我们预期标准普尔 500 收益率将与前 4 个因素中的一个或者多个有关,但是市场择机因素是与前 4 个因素不相关的。因为市场择机因素和标准普尔 500 收益率是不同的,我们一般不能用由相对于标准普尔 500 计算出的 CPAM 贝塔的比例来预期市场择机因素的敏感性。

另一个重要的宏观经济模型是所罗门-史密斯-巴尼 (Salomon Smith Barney) 美国股票风险特性模型,或者简称为所罗门 RAM。

例 11-16 所罗门 RAM 中的期望收益率

所罗门 RAM 模型用 9 个因素 (6 个宏观经济因素,1 个市场因素的残差,1 个规模因素的残差以及 1 个行业因素的残差) 来解释美国股票的收益率:

1. 经济增长率:工业产出的月度变化。
2. 信用水平:在去除与 30 年国库券收益率月度变动相关部分之后,所罗门-史密斯-巴尼高收益率市场 10 年以上指数的月度变动。
3. 长期利率:30 年国库券收益率的月度变动。
4. 短期利率:3 个月国库券收益率的月度变动。
5. 通货膨胀冲击:消费者价格指数月度变化中不可预期的部分。
6. 美元:加权交易美元的月度变化。
7. 市场残差:在去除上述 6 个因素效应之后的标准普尔 500 的月度收益率。
8. 小盘股溢价:在去除上述 7 个因素效应之后的罗素 2000 指数收益率和标准普尔 500 指数收益率的月度差别。
9. 行业残差:在去除前面 8 个因素效应之后的行业成分指数的月度收益率。

关于该模型的一些值得注意的要点如下:[⊖]

- 与 BIRR 模型和更一般的模型 (式 (11-16)) 不同,除了通货膨胀,其他所有的因素都是以收益率而不是以意外值来表示的。
- 因素 7、8 和 9 试图通过去除因素收益率中与之前因素相关的部分来获得净的或者由该因素所唯一贡献。因素 7、8 和 9 每个都是与其他 2 个以及其余因素不相关的;它们被称为正交 (orthogonal) (不相关) 的因素。另外,信用水平因素被设定为与长期利率不相关。
- 基于模型的解释力,每只股票接受了反映其决定系数的,从最高为 1 到最低为 5 的 RAM 排序。
- 因素敏感性是以标准化的形式给出的。

在表 11-17 中,因素敏感性都经过式 (11-19) 的标准化,以相同的单位进行了表示。

⊖ 关于这个模型的更多内容,参见 Sorenson, Samak 和 Miller (1998)。

表 11-17 4 只股票的因素敏感性

	银行	有线电视提供商	百货商店	计算机制造商
经济增长率	1.48	-1.30	-0.88	0.88
信用水平	-1.62	-0.77	-3.19	0.12
长期利率	-0.01	1.65	1.83	0.78
短期利率	1.00	-1.53	1.10	1.17
通货膨胀冲击	-0.97	-1.99	-1.57	0.80
美元	0.82	1.47	0.87	-1.63
市场残差	0.06	0.30	0.13	0.26
小盘股溢价	0.13	1.44	0.92	-1.19
行业残差	0.02	0.00	-0.03	0.00

注：每个格子中的数据都是标准化后的因素敏感性。

银行和计算机制造商发行投资级的债券。有线电视提供商和百货商店相对多地发行低于投资级的债券。计算机制造商很少在其资本结构中使用债务。只基于上述给定信息，回答下列问题。

- (1) 将有线电视提供商股票与那些平均股票因素敏感性进行比较。
- (2) 在其他条件相同的情况下，说明哪只股票在经济快速增长以及信用环境改善的经济中期望表现更好。
- (3) 提出一个可能的商业原因或者与公司基本面有关的原因，来解释有线电视提供商股票对于信用水平因素的敏感性为负的情况。

解：

- (1) 对于短期利率、交易加权的美元、市场残差因素和小盘股溢价因素的敏感性都在平均水平以上，因为它们表现出正的因素敏感性。对于行业残差因素的敏感性处于平均水平，因为它的因素敏感性为零。相反，有线电视提供商股票具有低于经济增长率、信用水平、长期利率和通货膨胀因素的平均敏感性，因为它表现出负的因素敏感性。
- (2) 银行和计算机制造商对于经济增长因素（工业产值的月度变动）具有正的敏感性。正如因素和因素敏感性所定义的，正的敏感性意味着在高速经济增长的环境中，当其他条件相同的情况下，这 2 只股票的收益率要高于平均水平。银行具有负的信用水平因素的系数，而计算机制造商具有正的敏感性。一个正在改善的信用环境意味着高收益债券的收益率正在下降。因此，我们会观察到在那种环境中的信用水平的值为负。在这 2 只股票中，我们预计只有具有负敏感性的银行股会在信用改善的环境中给予高于平均的收益率。因此，银行股预计在所述的情形中表现更好。
- (3) 信用水平因素本质上度量的是包含违约风险的溢价的变动。信用水平因素的系数为负意味着该股票在包含违约风险的溢价下降时（一个信用改善的环境）表现更好。对于有线电视提供商股票相对于信用水平因素的敏感性为负的一种解释是该公司非常喜欢以发行低于投资级债券的方式借款。这类债务的成本反映了明显的违约溢价。有线电视提供商的借款成本在信用改善的环境中应该会下降；该下降会对于其股价有一个正向的影响。

11.4.5.2 基本面因素模型

金融分析师经常因为各种各样的目的（包括投资组合业绩归因以及风险分析）来使用基本面因素模型。[Ⓓ] 基本面因素模型关注于使用观察到的描述证券本身或者证券发行者某种特性的基本面因素来解释个股的收益率。是否属于某个行业、市盈率、账面价值与股价比以及财务杠杆都是常见的基本面因素。

Ⓓ 投资组合业绩归因（portfolio performance attribution）从不同风险的来源对于投资组合的业绩表现进行了分析。

例 11-17 给出了一个对于宏观经济、基本面以及统计因素模型的研究。

例 11-17 其他的因素模型

康纳 (Connor) (1995) 为了比较哪个模型能够更好地解释股票的收益率, 而将一个宏观经济因素模型与一个基本面因素模型进行了对比。[Ⓐ]

康纳给出了对于 1985 年 1 月到 1993 年 12 月 779 只美国大盘股的月度收益率应用宏观经济因素模型的分析结果。通过使用 5 个宏观经济因素, 康纳得以近似解释 11% 的这些股票收益率的方差。[Ⓑ]表 11-18 给出了他的结果。

表 11-18 宏观经济因素的解释力

因素	只使用每种因素的解释力	将每种因素增加到其他因素之后解释力的增加
通货膨胀率	1.3%	0.0%
期限结构	1.1%	7.7%
工业产值	0.5%	0.3%
违约溢价	2.4%	8.1%
失业率	-0.3%	0.1%
所有因素		10.9%

资料来源: 康纳 (1995)。

康纳同样也给出了宏观经济因素分析中相同公司在基本面因素分析中的结果。所采用的因素模型为 BARRA US - E2 模型 (当前的版本为 E3)。表 11-19 给出了这些结果。在表中, “市场的变动” 表示股票的波动率, “成功” 是一个反映价格惯性的变量, “交易活动” 通过交易的频繁程度对于股票进行了区分, 而 “增长率” 通过过去和预期盈利增长率对股票进行了区分。[Ⓒ]

表 11-19 基本面因素的解释力

因素	只使用每种因素的解释力	将每种因素增加到其他因素之后解释力的增加
行业	16.3%	18.0%
市场的变动	4.3%	0.9%
成功	2.8%	0.8%
规模	1.4%	0.6%
交易活动	1.4%	0.5%
增长率	3.0%	0.4%
市盈率	2.2%	0.6%
账面价值与股价之比	1.5%	0.6%
盈利变动	2.5%	0.4%
财务杠杆	0.9%	0.5%
外国投资	0.7%	0.4%
劳动密集度	2.2%	0.5%
红利收益率	2.9%	0.4%
所有因素		42.6%

资料来源: 康纳 (1995)。

正如表 11-19 所示, 最重要的基本面因素是表示为 55 个虚拟变量的 “行业” 因素。相对于宏观经济因素模型大约 11% 的解释力, 基本面因素模型解释了大约 43% 的股票收益率的变动。康纳的论文没有提供不同因素在每个模型中统计检验的显著性; 然而, 康纳确实发现了基

Ⓐ 我们没有讨论的统计因素模型同样也在 Connor (1995) 中进行了报告。
Ⓑ 一个给定模型的解释力计算为 $1 - [(\text{所有股票收益率的特定资产方差的平均数}) / (\text{所有股票收益率总方差的平均数})]$ 。因为方差的估计值都对于自由度进行了修正, 所以一个因素对于解释力的边际贡献可能是 0 或者为负。解释力反映了所能解释的部分占给定模型所能解释平均股票水平收益率总方差的一个比例。

证据, 而该证据反映了在投资界大量使用这些模型的情况。例如, 基本面因素模型常被用于投资组合的业绩归因。我们将在本章的后面给出这方面应用的例子。通常, 基本面因素模型会比宏观经济模型使用更多的因素, 所以该模型也能给出更加详细的投资组合经理业绩结果的来源。

但是, 我们不能从这项研究中得出基本面因素模型自然地优于宏观经济因素模型的结论。每种主要类型的模型都有其应用的价值。在不同宏观经济因素模型中的每个因素都是由它们所代表的系统性风险 (即不能分散化的风险) 的统计证据所支持的。相反, 一位投资组合经理可以非常容易地排除一个特定行业来构建投资组合, 从而受到特定行业非系统性风险的影响。这两种类型的因素——宏观经济因素和基本面因素, 通常具有不同的度量和风险管理的方法。宏观经济因素集是很少的 (5 个变量)。基本面因素集是很多的 (67 个变量, 包括 55 个行业虚拟变量), 而以相对复杂为代价, 基本面因素模型可以用相对直接的与公司和证券相关的特征来给出风险更加详细的描述。康纳发现将宏观经济因素加入基本面因素模型之后, 它们没有边际解释力, 这意味着基本面风险特征完全刻画了用宏观经济因素贝塔所反映的风险特征。然而, 因为基本面因素提供了股票和其发行者非常详细的特征描述, 所以该结果并不令人惊讶。

我们已经了解了目前在业界实际应用的基本面模型中表示风险的不同因素。不仅在因素的特性和确切的定义上存在着不同, 同样在函数的形式以及股价的方法上也存在着许多不同。尽管存在着诸多不同, 但是我们可以将大部分基本面因素模型中刻画股票的因素归为下面 3 个大类:

- 公司基本面因素 (company fundamental factors)。这些因素与公司内部的业绩表现有关。例如与盈利增长有关的因素、与盈利变动有关的因素、与盈利惯性有关的因素以及与财务杠杆有关的因素。
- 公司股价关联因素 (company share-related factors)。这些因素包括估值指标以及其他与公司股票价格或者股票交易特征有关的其他指标。与前一类因素不同, 这些指标直接反映了投资者对于公司的期望。这些因素包括价格乘子如盈利收益率、红利收益率以及账面市值比。市值也属于该类因素。许多模型会使用的与股票价格惯性有关的变量, 股价波动率以及交易活动也都属于这类因素。
- 宏观经济因素 (macroeconomic factors)。行业类别或者反映是否处于该行业的指标就属于这类因素。许多模型中包含的诸如 CAPM 贝塔, 其他的反映系统性风险的类似指标以及收益率曲线的敏感性水平, 所有这些指标都属于这类因素。

11.4.6 应用

这一节我们会给出在投资实务中多因素模型的一些主要应用。

我们首先从对于投资组合业绩归因和风险分析开始进行讨论。我们可以用原始的收益率或者相对于投资组合基准的收益率来进行我们的讨论。因为它们提供了风险和收益率的一个参考标准, 所以基准在许多机构投资者对于控制风险的定量收益率的计划中发挥着重要的作用。因此, 我们应该关注于分析相对于基准的收益率。

多因素模型也能帮助投资组合经理从投资组合中确定所期望的风险特征。在讨论了业绩归因和风险分析之后, 我们将解释在构建具有类似于其他投资组合风险暴露的投资组合中多因素模型的应用。

11.4.6.1 分析收益率的来源

多因素模型可以帮助我们详细了解一位经理相对于基准的收益率的来源。为了简单起见, 在

本节中,我们分析一个完全投资于单一国内股票市场的股票投资组合的收益率的来源。^①

分析师经常在分解收益率的来源(分成基本元素)时,喜欢用基本面多因素模型。与统计因素模型不同,基本面因素模型允许用名称来描述投资组合业绩表现的来源。同样,与宏观经济因素模型不同,基本面模型更加直接也经常更加具体地给出投资风格选择和证券特征的建议。

我们首先需要理解积极管理经理的目标。正如前面所述的,经理通常相对于一个特定的基准来进行评估。积极投资组合经理持有不同于基准权重的证券,并试图通过这样的选择来增加投资组合相对于被动投资方法的价值。持有不同于基准权重的证券反映了投资组合经理不同于市场共同期望的预期。对于一位股票投资经理,这些预期可能与驱动股票收益率的共同因素有关,或者与公司的特性有关。因此,当我们评价一位积极投资经理时,我们想要问的问题是“该经理是否具有发现比被动策略收益更高的策略的洞察力?”使用多因素模型分析收益率的来源可以帮助我们回答这类问题。

一个投资组合 R_p 的收益率可以看作基准收益率 R_B 和积极收益率(active return)(投资组合收益率减去基准收益率)之和:

$$\text{积极收益率} = R_p - R_B \quad (11-20)$$

有了一个因素模型之后,我们就可以分析投资组合经理的作为两部分之和的积极收益率。第1部分是投资组合经理因素乘子(积极因素的敏感性)和因素收益率的乘积;我们可以称该部分为由因素乘子带来的收益率。第2部分是反映经理单个资产选择能力的积极收益率部分;我们可以称该部分为资产选择(asset selection)。式(11-21)给出了积极收益率所分解出的2部分:

$$\text{积极收益率} = \sum_{j=1}^K [(投资组合敏感性)_j - (基准的敏感性)_j] \times (因素收益率)_j + \text{资产选择} \quad (11-21)$$

在式(11-21)中,我们度量了评估期初投资组合相对于我们风险模型中每个因素的敏感性和基准相对于我们风险模型中每个因素的敏感性。

例11-18举例说明了一个相对简单的基本面因素模型在分解和解释收益率中的使用。

例11-18 一位股票投资组合经理积极收益率的分解

作为一位养老基金出资人的股票分析师,罗纳德服务(Ronald Service)使用下面的多因素模型来评价美国股票投资组合:

$$R_p - R_F = a_p + b_{p1}RMRF + b_{p2}SMB + b_{p3}HML + b_{p4}WML + \epsilon_p \quad (11-22)$$

式中, R_p 和 R_F 分别表示投资组合的收益率和无风险收益率; RMRF 表示市值加权股票指数收益率超出一个月国库券的收益率的部分; SMB 表示小减去大,它是一个规模(市值)因素。SMB 是3个小市值投资组合的平均收益率减去3个大市值投资组合平均收益率的值; HML 表示高减去低,2个最高账面市值比的投资组合的平均收益率减去2个最低账面市值比的投资组合的平均收益率; WML 表示赢家减去输家,它是一个惯性因素。WML 是过去一年赢家的投资组合收益率减去过去一年输家的投资组合收益率。^②

① 该假设允许我们忽略国家选择、资产配置、市场择机以及货币对冲的影响,从而极大地简化了分析。然而,甚至在一个更加一般的情况下,我们仍旧可以使用多因素模型进行类似的分析。

② WML 是一个滞后1个月的11个月收益率最高30%部分的等权重加权平均数减去滞后1个月的11个月收益率最低30%部分的等权重加权平均数。该模型是基于 Carhart (1997) 得到的; WML 是 Carhart 中的 PRIYR 因素。

在式 (11-22) 中, 敏感性被解释为回归系数, 并且没有经过标准化。

罗纳德服务的当前工作是评估大部分目前雇佣的美国股票投资经理的业绩表现。这些经理的基准是罗素 1000 指数, 一个反映美国大盘股表现的指数。经理将其自己描述为一位“选股者”, 并指出她击败基准的表现就是其成功的证据。表 11-20 给出了基于反映经理在该年积极收益率来源的式 (11-21) 的分析。在表 11-20 中, “A. 从因素乘子所获得的收益率”, 等于 2.1241%, 等于该栏上面 4 个数的总和; 从因素乘子所获得的收益率是式 (11-21) 的第 1 部分。表 11-20 给出了资产选择等于 -0.05%。积极收益率为 $2.1241\% + (-0.05\%) = 2.0741\%$ 。

表 11-20 积极收益率的分解

因素	因素敏感性			对于积极收益率的贡献		
	投资组合 (1)	基准 (2)	差别 (3) = (2) - (1)	因素收益率 (4)	绝对值 (3) × (4)	占总积极收益率的比例
RMRF	0.85	0.90	-0.05	5.52%	-0.276 0	-13.3%
SMB	0.05	0.10	-0.05	-3.35%	0.167 5	8.1%
HML	1.40	1.00	0.40	5.10%	2.040 0	98.4%
WML	0.08	0.06	0.02	9.63%	0.192 6	9.3%
A. 来自因素乘子的收益率					2.124 1	102.4%
B. 资产选择					-0.050 0	-2.4%
C. 积极收益率 (A + B)					2.074 1	100.0%

从他先前的工作中, 罗纳德服务了解到在式 (11-22) 中成长风格投资组合的收益率经常具有正的惯性因素 (WML) 敏感性。相反, 对于价值风格投资组合的收益率, 例如采用反转策略的投资组合, 通常具有很低的或者负的惯性因素的敏感性。

使用这些信息, 回答下列问题:

- (1) 确定投资经理的投资风格。
- (2) 评估投资经理一年积极收益率的来源。
- (3) 对于收益率分解的结果, 罗纳德服务可能会讨论关于投资经理的什么问题?

解:

(1) 基准的敏感性反映了一位投资经理所选投资机会集基准的风险特征。我们可以通过检验其所选择基准的特征来推断投资经理预期的风格。于是我们通过检验投资组合实际因素的风险暴露来确认这些推断:

- 具有高账面市值比的股票一般被认为是价值股。因为养老基金出资人选择的基准具有对于 HML (高账面市值比因素减去低账面市值比因素) 高的敏感性 (1.0), 所以我们能够推断投资经理具有价值型的投资风格。实际 HML 的敏感性为 1.4 说明投资经理比基准具有更多高账面市值比的风险暴露。
- 基准和投资组合相对于惯性因素近似中性是与价值型的投资风格相一致的。
- 基准和投资组合相对于 SMB 低的风险暴露本质上说明其对于小盘股是没有净风险暴露的。

综合上述的所有考虑, 由此说明该投资经理具有一个大盘股价值型的投资风格。

(2) 投资经理正的积极收益率的最重要来源是她相对于 HML 因素正的积极风险暴露。该因素贡献了 $(1.40 - 1.00)(5.01\%) = 2.04\%$ 或者近似 98% 的已实现的约为 2.07% 的积极收益率。在评估期中, 投资经理加强了价值型的投资风格, 并且该风格带来了回报。投资经理对于

整个市场 (RMRF) 的积极风险暴露是没有获利的, 但是她对于小盘股 (SMB) 和惯性 (WML) 的积极风险暴露是获利的; 然而, 投资经理相对于 RMRF、SMB 和 WML 的积极风险暴露规模是相对较小的, 因此, 这些积极收益率因素的影响相对于 HML 巨大而成功的影响是微不足道的。

(3) 虽然投资经理描述自己是一个“股票选择者”, 但是她在这个时间段资产选择的积极收益率却为负。她正的积极收益率是源于对于 HML 的很大的积极风险暴露以及该因素在这个时间段的高收益率。如果市场更有利于成长股而不是价值股并且投资经理在挑选各个资产方面没有做得更好, 那么该投资经理的业绩表现将是无法令人满意的。投资经理能否提供她能够预测 HML 因素收益率变化的证据? 她是否对于其择股能力过于自信? 罗纳德服务可能想要讨论这些关于投资经理的内容。收益率的分解已经帮助其区分了由价值型与成长型的投资风格带来的收益率以及由投资经理选择的投资方法所挑选的个股带来的收益率。

11.4.6.2 分析风险的来源

我们继续对于积极收益率进行关注, 在这一节中我们对于积极风险进行分析。积极风险 (active risk) 是指积极收益率的标准差。有许多与此表示完全相同概念的术语, 因此, 我们需要简要地来了解一下它们。一个经常使用的同义词是跟踪误差 (tracking error, TE), 但是除非读者所理解的误差是指标准差, 不然该术语就可能会产生误解; 跟踪误差的波动率 (tracking-error volatility, TEV) 也被人们使用 (这个误差被理解为一种差别); 现在, 跟踪风险 (tracking risk) 也被广泛使用 (但是 TR 的缩写可能被误解为总收益率)。这里, 我们将使用缩写 TE 来表示跟踪风险的概念, 并且我们将跟踪风险表示为:

$$TE = s(R_p - R_B) \quad (11-23)$$

在式 (11-23) 中, $s(R_p - R_B)$ 表示我们采用投资组合收益率 R_p 和基准收益率 R_B 之差的时间序列的样本标准差 (以 s 表示)。我们应该注意的是积极收益率和跟踪风险是以相同的时间基准进行表示的。^①

作为一个广泛使用的跟踪风险变化范围的指标, 在美国股票市场中一个良好的被动投资策略经常能够控制跟踪风险在每年 1% 以内。一个半积极或者增强型的指数投资策略, 其在很大程度上利用投资经理的预期, 通常具有的跟踪风险目标为每年 2%。一个分散化的积极型美国大盘股策略, 其可能的基准为标准普尔 500 指数, 通常具有的跟踪风险的范围为每年 4% ~ 6%。一个激进的积极型股票投资经理可能具有的跟踪风险的范围为每年 6% ~ 9% 或者更多。

与我们使用传统的夏普指标来评价绝对收益率有些类似, 平均积极收益率和积极风险之比, 即信息比率 (information ration, IR), 是一个用来评价每单位积极风险的平均积极收益率的工具。历史的或者事后的 IR 具有如下形式:

$$IR = \frac{\bar{R}_p - \bar{R}_B}{s(R_p - R_B)} \quad (11-24)$$

式 (11-24) 的分子, $\bar{R}_p$ 和 $\bar{R}_B$ 分别代表投资组合的样本平均收益率和基准的样本平均收益率。为了举例说明其计算方法, 我们假设一个投资组合实现了 9% 的平均收益率, 而在相同的时间段其基准获得了一个 7.5% 的平均收益率, 而投资组合的跟踪风险是 6%, 那么我们可以计算得到其

① 为了将每天基于日收益率的 TE 年化, 我们基于一年 250 个交易日将每天的 TE 乘以 $(250)^{1/2}$; 为了将每月基于月收益率的 TE 进行年化, 我们将月度 TE 乘以 $(12)^{1/2}$ 。

信息比率为 $(9\% - 7.5\%) / 6\% = 0.25$ 。设定可接受的积极风险或者跟踪风险的标准是一些机构投资者试图保证他们投资的总体风险和风格特征与他们的期望相一致的一种方法。

例 11-19 与投资经理交流

积极收益率和积极风险的框架对于那些想要密切控制投资风险的投资者很有吸引力。基准作为一个已知的、连续观测的与所给出的和传达的定量风险和收益率目标有关的参考标准。例如，一个美国公共员工退休系统给预期未来的投资经理发出一个关于“控制风险的大盘股票基金”请求（或者对于计划的要求），该请求包括如下要求：

- 所选的股票必须为标准普尔 500 的成分股。
- 投资组合应该最少包含 200 只股票。在购买时，投资任一股票的最大金额为投资组合市场价值的 5% 和股票在标准普尔 500 指数中所占权重的 150% 这 2 个数中的较大者。
- 或者自初始发行起或者在最后 7 年中，投资组合必须具有最小 0.30% 的信息比率。
- 或者自初始发行起或者在最后 7 年中，投资组合必须同时具有低于 4% 的关于标准普尔 500 的跟踪风险。

分析师使用多因素模型来详细了解一位投资组合经理的风险暴露。在分解积极风险过程中，分析师的目标是度量投资组合的每个维度的积极风险暴露。换句话说，了解跟踪风险的来源。^① 分析师想要回答的问题有下面这些：

- 哪些积极风险暴露对于投资经理的跟踪风险贡献最大？
- 投资组合经理是否意识到其积极风险暴露的性质，如果是这样的话，那么他是否能够用一个清晰的原理来表示它们？
- 投资组合的积极风险暴露是否与投资经理所声称的投资理念相一致？
- 哪种积极乘子获得了充分的该积极风险水平下的收益率？

在回答这些问题时，分析师经常选择基本面因素模型，因为它们可以用来将积极风险暴露与投资经理投资组合决策以相当直接和直观的方式相联系。在本节中，我们将解释如何使用多因素模型来分解和解释一个投资组合的积极风险。

我们先前讨论了积极收益率的分解；现在我们将讨论积极风险的分解。在分析风险时，使用方差而不是标准差是非常方便的，因为不相关变量的方差是可加的。我们将积极风险的方差称为积极风险平方（active risk squared）：

$$\text{积极风险平方} = s^2(R_p - R_B) \quad (11-25)$$

我们可以将一个投资组合的积极风险平方分成 2 部分：

- **积极因素风险**（active factor risk）是指由风险模型中具体设定的因素所带来的、超出基准（水平）部分的风险暴露对于积极风险平方的贡献。^②

① 投资组合的积极风险是成分股积极风险的加权平均值。因此，我们也可以在单个持有的情况下进行这一分析。因为投资组合经理可能发现这种方法在对其积极风险投资进行调整时非常有用。

② 在整个讨论中，“积极”意味着“不同于基准”。

- 积极特有风险 (active specific risk) 或者资产选择风险 (asset selection risk) 是指投资组合中各个资产按积极权重加权所得部分对于积极风险平方的贡献, 其权重是与资产残差风险相联系的。[Ⓐ]

当我们将积极风险平方应用到单个资产类中时, 积极风险平方就具有两个部分。积极风险平方两部分的分解为:

积极风险平方 = 积极因素风险 + 积极特有风险 (11-26)

积极因素风险代表积极风险平方能够被投资组合积极因素风险暴露所解释的部分。积极因素风险可以间接地表示为积极风险平方和积极特有风险之差, 而积极特有风险的表达式为[Ⓑ]

积极特有风险 = $\sum_{i=1}^n w_i^a \sigma_{\epsilon_i}^2$

式中, w_i^a 表示第 i 个资产在投资组合中的积极权重 (即投资组合中资产的权重与基准中其权重的差), 而 $\sigma_{\epsilon_i}^2$ 表示第 i 个资产的残差风险 (第 i 个资产收益率不能被因素所解释部分的方差)。[Ⓒ]

积极特有风险找出了投资经理所承担的无法被因素所解释的积极风险或者说是残差风险。我们应该寻找一个正的资产选择的平均收益率来作为承担积极特有风险的补偿。

例 11-20 积极风险的一个比较

理查德·格雷 (Richard Gray) 正在比较 4 位具有相同基准的美国股票基金经理的投资风险。他使用一个基本面因素模型——BARRA US-E3 模型, 其中包含了 13 个风险指数和一组 52 个行业类别。风险指数度量了公司及其股票的不同基本面情况, 如公司规模、财务杠杆以及红利收益率。在模型中, 公司对于所有其经营的行业具有非零的风险暴露。表 11-21 基于式 (11-26), 给出了格雷对于这 4 位经理的积极风险平方的分析。[Ⓓ]在表 11-21 中, 标记为“行业”的一栏给出了投资组合与其持有的行业风险暴露有关的积极因素风险; “风险指数”一栏给出了投资组合与其持有的 13 个风险指数有关的积极因素风险。

表 11-21 积极风险平方的分解

投资组合	积极因素			积极特有风险	积极风险平方
	行业	风险指数	总因素		
A	12.25	17.15	29.40	19.60	49
B	1.25	13.75	15.00	10.00	25
C	1.25	17.50	18.75	6.25	25
D	0.03	0.47	0.50	0.50	1

注: 每一格的单位都是百分比的平方。

Ⓐ 正如我们所使用的术语, “积极特有风险”和“积极因素风险”用的是方差而不是标准差。
Ⓑ 直接的计算积极因素风险的方法如下。一个投资组合对于一个给定的因素 j , 其积极因素风险暴露可以通过将每个资产对于因素 j 的敏感性用其积极权重进行加权, 然后加总各项 $b_i^a = \sum_{i=1}^n w_i^a b_{ji}$ 。然后积极因素风险等于 $\sum_{i=1}^K \sum_{j=1}^K b_i^a b_j^a \text{Cov}(F_i, F_j)$ 。
Ⓒ 资产收益率的残差被认为不仅相互间是不相关的, 而且与其他因素收益率也不相关。
Ⓓ 在积极因素风险中存在一个协方差的概念, 其反映行业成员与风险指数的相关系数。但是在这个例子中我们假设它是可以忽略的。

使用表 11-21 中给出的信息，回答下列问题：

- (1) 比较投资组合 A 和 B 积极风险的分解情况。
- (2) 比较投资组合 B 和 C 积极风险的分解情况。
- (3) 描述投资组合 D 的投资方式。

解：

(1) 表 11-22 再次给出了表 11-21 的信息，以显示出不同积极风险来源贡献的比例。在表 11-22 的最后一栏中，我们现在给出的是积极风险平方的平方根，即积极风险或者跟踪风险。为了解释表 11-22 中间的几栏，投资组合 A 在行业一栏下的值为 25%，其可以由 $12.25/49 = 0.25$ 得到。因此投资组合 A 与行业风险暴露相关的积极风险为积极风险平方的 25%。

表 11-22 积极风险的分解（重复的）

投资组合	积极因素 (占总积极风险的比例(%))			积极特有风险 (占总积极风险的比例(%))	积极风险平方
	行业	风险指数	总因素		
A	25	35	60	40	7
B	5	55	60	40	5
C	5	70	75	25	5
D	3	47	50	50	1

投资组合 A 假设比 B 的积极风险水平更高（跟踪风险 7% 比上 5%）。投资组合 A 和 B 假设具有相同比例的积极因素风险和积极特有风险，但是这两个投资组合在积极因素风险暴露的类型上存在显著的不同。投资组合 A 承担着很大的积极行业风险，而投资组合 B 相对于基准近似于行业中性。相反，投资组合 B 对于代表公司及其股票特征的风险指数具有很高的积极乘子。

(2) 投资组合 B 和 C 在它们积极风险的绝对量上是类似的。而且，投资组合 B 和 C 都近似相对于基准是行业中性的。投资组合 C 相对于风险指数承担更多的积极因素风险，但是投资组合 B 承担更多的积极特有风险。我们也可以从第 2 点推断出 B 的分散化程度稍稍不如 C。

(3) 从其可以几乎忽略的积极风险水平判断，投资组合 D 表现为被动管理的投资组合。由表 11-21，投资组合 D 的积极因素风险为 0.50，以标准差表示为 0.707%，其表明该投资组合在模型所识别出的驱动平均收益率的风险维度上与基准非常接近。

例 11-20 给出了一组具有不同跟踪风险的虚拟投资组合，其中的积极因素风险倾向于高于积极特有风险。给定一个良好设定的多因素模型和一个完全分散化的投资组合，该关系是相当普通的。对于完全分散化的投资组合，管理积极因素风险一般的主要工作是管理跟踪风险。例 11-20 给出了一个在总量水平上对于积极风险的分析；一个关于包含 13 个风险指数的多因素模型的积极因素风险被加总为一个数字。在评价业绩表现中，一位分析师可能会关心一个更加详尽的投资组合的积极风险。分析师是如何来评价一位投资经理积极因素风险暴露对于积极风险平方的单个贡献的呢？

无论因素是什么，评价积极因素风险暴露对于积极风险平方贡献的方法是一样的。这个量被称为因素对于积极风险平方的边际贡献（FMCAR）。在 K 个因素的情况下，因素 j 对于积极风险平方的边际贡献 $FMCAR_j$ 为

$$FMCAR_j = \frac{b_j^a \sum_{i=1}^K b_i^a Cov(F_j, F_i)}{\text{积极风险平方}}$$

(11-27)

其中 b_j^a 是投资组合相对于因素 j 的积极风险暴露。分子是因素 j 的积极因素风险。[⊖] 分子类似于我们先前在讨论市场模型中所遇到的贝塔的表达式,但是具有更多的因素,因素协方差和方法是相关的。为了在一个简单的情形下对于式(11-27)进行说明,假设我们有一个两因素模型:

- 投资经理对于第1个因素的积极风险暴露为0.50;即 $b_1^a = 0.50$ 。另一个积极因素风险暴露为 $b_2^a = 0.50$ 。
- 该因素的方差—协方差矩阵描述为 $\text{Cov}(F_1, F_1) = \sigma^2(F_1) = 225$, $\text{Cov}(F_1, F_2) = 12$, 而 $\text{Cov}(F_2, F_2) = \sigma^2(F_2) = 144$ 。
- 积极特有风险为53.71。

我们首先计算每个因素的积极因素风险;该计算表达式是式(11-27)的分子。那么我们通过加总积极因素风险和积极特有风险得到积极风险平方,并得到式(11-27)给出的比率。对于第1个因素,我们计算出的 FMCAR_1 的分子为

$$b_1^a \sum_{i=1}^2 b_i^a \text{Cov}(F_1, F_i) = 0.50[0.50(225) + 0.15(12)] = 57.15$$

对于第2个因素,我们有

$$b_2^a \sum_{i=1}^2 b_i^a \text{Cov}(F_2, F_i) = 0.15[0.50(12) + 0.15(144)] = 4.14$$

积极因素风险为 $57.15 + 4.14 = 61.29$ 。加上积极特有风险,我们求出积极风险平方为 $61.29 + 53.71 = 115$ 。因此,我们有 $\text{FMCAR}_1 = 57.15/115 = 0.497$ 或者 49.7%, 而 $\text{FMCAR}_2 = 4.14/115 = 0.036$ 或者 3.6%。积极因素风险占总风险中的比例为 $\text{FMCAR}_1 + \text{FMCAR}_2 = 49.7\% + 3.6\% = 53.3\%$ 。积极特有风险贡献了积极风险平方的 $100\% - 53.3\% = 46.7\%$ 。例11-21举例说明了这些概念的应用。

例 11-21 各个积极因素风险的分析

威廉·维泽尔(William Whetzell)负责月度组织内部的业绩归因以及由其组织(加拿大捐赠基金会)管理的国内核心股票型基金的风险分析。在他的月度分析中,维泽尔使用包含下列因素的风险模型:

- 市值的对数
- E/P, 盈利收益率
- B/P, 账面市值比
- 盈利增长率
- 平均红利收益率
- D/A, 长期负债资产比率
- 股权收益率(ROE)的波动率
- EPS的波动率

该模型的因素敏感性具有基本面因素模型中因素敏感性的标准解释。

⊖ 如果我们将分子从 $j=1$ 加到 K , 那么我们将得到脚注74所给出的积极因素风险的表达式。

在确定了他在最后一个财政年度获得了 0.75% 的积极收益率之后，维泽尔转而对于投资组合的风险进行分析。

在该财政年度之初，投资委员会具有下列决策：

- 决定战术性地增加投资组合中小盘股的比重；
- 在证券选择时采用在“合理价位的盈利增长”（earnings growth at a reasonable price, GARP）原则；
- 保持任何积极因素风险在以标准差表示的每年 5% 的水平；
- 保持积极特有风险不超过积极风险平方的 50%；
- 使得信息比率大于等于 0.15。

在维泽尔给出其报告之前，一位投资委员会的成员审阅了他的资料，并建议投资委员会如果股票型基金持续保持最后一个财政年度表现的情况下，就应该对国内股票采用一个被动投资策略。

表 11-23 给出了股票型基金的信息。所构建的因素收益率是近似相互不相关的。

表 11-23 风险分析数据

敏感性			
因素	投资组合	基准	因素的方差
市值的对数	0.05	0.25	225
E/P	-0.05	0.05	144
B/P	-0.25	-0.02	100
盈利增长率	0.25	0.10	196
红利收益率	0.01	0.00	169
D/A	0.03	0.03	81
ROE 的波动率	-0.25	0.02	121
EPS 的波动率	-0.10	0.03	64
积极特有风险 =			29.940 6
积极特定收益率 =			-0.5%
积极收益率 =			0.75%

基于表 11-23 中的信息，回答下列问题：

- (1) 对于每个因素，计算 (A) 积极因素风险，(B) 积极风险平方的边际贡献。
- (2) 讨论是否数据与投资委员会所要满足的目标相一致。
- (3) 评估该年捐赠者的风险调整后的业绩。
- (4) 给出 2 条支持的委员会成员关于被动投资策略建议的证据。

解：

(1)

(A) 一个因素的积极因素风险 = (因素的积极敏感性)² (因素的方差)，在式 (11-27) 中没有因素的相关系数。

- 市值的对数 = $(0.05 - 0.25)^2 (225) = 9.0$
- $E/P = (0.05 - 0.05)^2 (144) = 1.44$
- $B/P = [-0.25 - (-0.02)]^2 (100) = 5.29$
- 盈利增长率 = $(0.25 - 0.10)^2 (196) = 4.41$

- 红利收益率 = $(0.01 - 0.00)^2(169) = 0.0169$
- $D/A = (0.03 - 0.03)^2(81) = 0.0$
- ROE 的波动率 = $(-0.25 - 0.02)^2(121) = 8.8209$
- EPS 的波动率 = $(-0.10 - 0.03)^2(64) = 1.0816$

(B) 各个积极因素风险之和为 30.0594。我们在该和的基础上加入积极特有风险，获得积极风险平方为 $30.0594 + 29.9406 = 60$ 。因此，各个因素的 FMCAR 如下：

- 市值的对数 = $9/60 = 0.15$
- $E/P = 1.44/60 = 0.024$
- $B/P = 5.29/60 = 0.0882$
- 盈利增长率 = $4.41/60 = 0.0735$
- 红利收益率 = $0.0169/60 = 0.0003$
- $D/A = 0.0/60 = 0.0$
- ROE 的波动率 = $8.8209/60 = 0.1470$
- EPS 的波动率 = $1.0816/60 = 0.0180$

(2) 我们依次来考虑每个投资委员会的目标。第1个目标是战术性地提高投资组合中小盘股的比重。对于市值对数因素敏感性为0意味着股票规模的平均风险暴露。给定基本面因素模型中因素敏感性的标准解释，该风险暴露为1意味着增加规模会带来一个对于收益率正的风险暴露，其高于均值一个标准差。虽然股票型基金相对于规模的风险暴露为正，但是积极风险暴露为负。这个结果与增加小盘股的权重的论述是一致的。

第2个目标是在证券选择时采用在“合理价位的盈利增长”(earnings growth at a reasonable price, GARP)原则。股票型基金相对于盈利增长率具有正的积极风险暴露，其与所选择的公司具有高盈利增长率是一致的。然而问题是该方法中的“合理价格”部分是否会被满足。该基金的E/P和B/P敏感性为负，表明低于平均水平的盈利收益率和B/P（高于平均水平的P/E和P/B）。这些因素的积极风险暴露也为负。如果高于平均水平的盈利增长率在市场中被定价，那么该基金可能需要持有负的积极风险暴露，因此我们不能获得最终的结论。所以我们可以认为没有数据正面支持执行GARP策略的结论。

第3个目标是保持任何积极因素风险在以标准差表示的每年5%的水平。最大的积极因素风险是市值对数的因素。以一个标准差表示，该风险为每年 $(9)^{1/2} = 3\%$ 。因此，该目标是得到满足的。

第4个目标是保持积极特有风险不超过积极风险平方的50%。积极特有风险占积极风险平方的 $29.9406/60 = 0.4990$ 或者 49.9%，因此，该目标得到满足。

第5个目标是要使得信息比率大于等于0.15。信息比率是积极收益率除以积极方差（跟踪方差）。已知积极收益率为0.75%。积极风险是积极风险平方的平方根： $(60)^{1/2} = 7.7460\%$ 。因此， $IR = 0.75\% / 7.7460\% = 0.0968$ 或者近似为0.10，其低于所给出的目标0.15。

(3) 捐赠基金实现的信息比率为0.10，这意味着该年其风险调整后的业绩表现是不令人满意的。

(4) 虽然基金承担了很大的积极特有风险，但是该年的积极特定收益率为负。因此，特定风险具有负的收益。而且实现的信息比率低于投资委员会的目标。根据这些只基于一年的分析，我们认为这些事实说明了成本有效的不同于指数化的策略。

在我们对于业绩归因和风险分析的讨论中，我们已经给出了关于普通股的例子。多因素模型同样也在债券或者其他资产类的投资组合应用中发挥着同样的作用。

我们已经举例说明了在分析一个投资组合积极收益率和积极风险中多因素模型的应用。至少在投资组合构建中使用多因素模型也是同样重要的。在投资组合管理过程这个阶段，多因素模型允许投资组合经理聚焦于相对于其基准风险的判断或者管理其相对于基准风险的投资组合风险控制。在余下的章节中，我们将讨论多因素模型的这些应用。

11.4.6.3 因素投资组合

一位投资组合经理可以使用多因素模型来建立一个反映特定风险类的投资组合。例如，他可能想要建立并使用一个因素投资组合。对于某个因素的因素投资组合是指对于该因素的敏感性为1，而对于其他因素的敏感性都是0的投资组合。因此，该投资组合至对于一个风险因素存在风险暴露，其能完全反映该因素的风险。因为完全反映该风险的来源，所以因素投资组合是想要对冲（抵消）该风险或者利用该风险投机的投资组合经理所关注。例 11-22 举例说明了因素投资组合的这一应用。

例 11-22 因素投资组合

分析师万达·史密斯菲尔德（Wanda Smithfield）已经构建了6个可能被其公司投资组合经理使用的投资组合。这些投资组合分别被标记为表 11-24 中的 A，B，C，D，E 和 F。

表 11-24 因素投资组合

风险因素	投资组合					
	A	B	C	D	E	F
信用风险	0.50	0.00	1.00	0.00	0.00	0.80
投资期限风险	1.92	0.00	1.00	1.00	1.00	1.00
通货膨胀风险	0.00	0.00	1.00	0.00	0.00	-1.05
商业周期风险	1.00	1.00	0.00	0.00	1.00	0.30
市场择机风险	0.90	0.00	1.00	0.00	0.00	0.75

注：每一格中都表示因素敏感性。

- (1) 投资组合经理想要对于未来实际商业活动增长进行猜测。
A. 确定给出的6个投资组合中对于该投资经理最有用的投资组合，并予以证明。
B. 该投资经理将会在问题 A 的投资组合中持有何种类型的头寸。
- (2) 投资组合经理想要对冲已知的投资期限为正的风险暴露。
A. 确定给出的6个投资组合中对于该投资经理最有用的投资组合，并予以证明。
B. 该投资经理将会在问题 A 的投资组合中持有何种类型的头寸。

解：

(1A) 投资组合 B 是最适合的选择。投资组合 B 是反映商业周期的因素投资组合，因为其对于商业周期风险的敏感性为1，而对于所有其他因素的敏感性为0。所以投资组合 B 是对于未来实际商业活动增长的猜测是有效的。

(1B) 投资组合经理将通过持有投资组合 B 的多头头寸来对于未来实际商业活动的增长进行猜测。

(2A) 投资组合 D 是一个合适的选择。投资组合是反映投资期限风险的因素投资组合，因为其对于投资期限风险的敏感性为1，而对于所有其他因素的敏感性为0。所以投资组合 D 对于对冲已知的投资期限正的风险暴露是有效的。

(2B) 投资组合经理将通过持有投资组合 D 的空头头寸来对冲已知的投资期限正的风险暴露。

下一节将给出构建具有所需因素敏感性配置投资组合的具体步骤。

11.4.6.4 构建一个跟踪投资组合

在前一节中，我们讨论了多因素模型在投机或者对冲特定因素风险中的应用。也许更加一般地，投资组合经理会使用多因素模型来控制投资组合相对于基准的风险。例如，在一个风险控制的积极或者增强的指数策略中，投资组合经理可能试图获得超过基准的小额增量收益率，而同时通过将其投资组合中的因素敏感性与基准进行匹配来控制风险。由此得到的投资组合将是跟踪投资组合的一个例子。一个跟踪投资组合（tracking portfolio）是指具有和基准或者其他投资组合相匹配的因素敏感性的投资组合。

构造具有一组目标因素敏感性的投资组合的技术涉及利用代数方法求解一组方程的解。

计算约束的数量。每个贝塔的目标值代表投资组合的一个约束，而另一个约束是投资组合中所有投资的权重之和必须等于 1。因为有约束条件，所以所有的投资都是需要的。

对于投资组合投资的权重建立一个方程来反映每个对于投资组合的约束。我们具有一个反映投资组合的权重加总为 1 的方程。我们对于每个目标因素敏感性具有一个方程；在等号的左边，我们具有投资因素敏感性的关于因素的加权平均，而在等号的右边我们具有目标因素敏感性。

通过求解方程组来解出投资组合投资的权重。

在例 11-23 中，我们举例说明了一个跟踪投资组合是如何被构建的。

例 11-23 构建一个跟踪投资组合

假如一个养老计划的出资人想要完全投资于美国普通股。该养老计划的出资人已经对于该投资组合经理设定了一个股票基准，而该投资组合经理决定对于该基准构建一个跟踪投资组合。为了使用熟悉数据的缘故，我们继续使用投资组合 J、K 和 L，以及例 11-13 中相同的两因素模型。

投资组合经理确定出的基准具有 1.3 的通货膨胀意外值的敏感性，以及 1.975 的 GDP 意外值的敏感性。一共有 3 个约束。第 1 个约束是投资组合的权重加总为 1，第 2 个约束是通货膨胀因素敏感性的加权之和为 1.3（与基准相匹配），而第 3 个约束是 GDP 因素敏感性的加权之和为 1.975（与基准相匹配）。因此，我们需要 3 项投资组合 J、K 和 L 来形成该投资组合。

表 11-14 （重复的）一个两因素模型的样本投资组合

投资组合	期望收益率	通货膨胀因素的敏感性	GDP 因素的敏感性
J	0.14	1.0	1.5
K	0.12	0.5	1.0
L	0.11	1.3	1.1

方程 1. 该方程反映了投资组合的权重之和必须等于 1。

$$\omega_J + \omega_K + \omega_L = 1$$

方程 2. 第 2 个方程反映了投资组合 J、K 和 L 通货膨胀意外值的敏感性加权平均值必须等于基准的通货膨胀意外值的敏感性 1.3。该要求保证了跟踪投资组合具有和基准一样的通货膨胀风险。

$$1.0\omega_J + 0.5\omega_K + 1.3\omega_L = 1.3$$

正如我们所提到的，我们需要 3 个方程来确定投资组合中投资组合权重 ω_J ， ω_K 和 ω_L 。

方程 3. 第 3 个方程反映了投资组合 J、K 和 L 的 GDP 意外值的敏感性加权平均值必须等于基准的 GDP 意外值的敏感性 1.975。该要求保证了跟踪投资组合具有和基准一样的 GDP 风险。

$$1.5\omega_J + 1.0\omega_K + 1.1\omega_L = 1.975$$

我们可以按如下方法求解出权重。从方程 1 中有 $\omega_L = (1 - \omega_J - \omega_K)$ 。我们将该结果代入另外 2 个方程之中得到

$$1.0\omega_J + 0.5\omega_K + 1.3(1 - \omega_J - \omega_K) = 1.3, \text{ 化简为 } \omega_K = -0.375\omega_J$$

以及

$$1.5\omega_J + 1.0\omega_K + 1.1(1 - \omega_J - \omega_K) = 1.975, \text{ 化简为 } 0.4\omega_J - 0.1\omega_K = 0.875$$

接着我们将 $\omega_K = -0.375\omega_J$ 代入 $0.4\omega_J - 0.1\omega_K = 0.875$, 得到 $0.4\omega_J - 0.1(-0.375\omega_J) = 0.875$ 或者 $0.4\omega_J + 0.0375\omega_J = 0.875$, 因此, $\omega_J = 2$ 。

利用之前得到的 $\omega_K = -0.375\omega_J$, $\omega_K = -0.375 \times 2 = -0.75$ 。最后, 从 $\omega_L = (1 - \omega_J - \omega_K) = [1 - 2 - (-0.75)] = -0.25$ 。最后我们得到,

$$\omega_J = 2$$

$$\omega_K = -0.75$$

$$\omega_L = -0.25$$

跟踪投资组合的期望收益率为 $0.14\omega_J + 0.12\omega_K + 0.11\omega_L = 0.14(2) + 0.12(-0.75) + 0.11(-0.25) = 0.28 - 0.09 - 0.0275 = 0.1625$ 。使用例 11-13 中相同的输入变量, 我们所计算出的 APT 模型为 $E(R_p) = 0.07 - 0.02\beta_{p,1} + 0.06\beta_{p,2}$ 。对于跟踪投资组合, $\beta_{p,1} = 1.3$, $\beta_{p,2} = 1.975$ 。因为 $E(R_p) = 0.07 - 0.02(1.3) + 0.06(1.975) = 0.1625$, 所以我们证实了期望收益率的计算。

11.4.7 总结

在前面的几节中, 我们说明了多因素模型是如何帮助投资组合经理解决实际度量和控制风险的实际工作的。我们现在要讨论 CAPM 和 APT 之间的不同之处, 并提供其他问题(为什么一些风险可以被定价以及多因素模型所描述的投资组合与 CAPM 描述的投资组合有什么不同)的一些解释。一个投资者可以通过使用多因素模型而不是单因素模型, 获得更好的投资组合决策。

CAPM 提供了投资者对于投资思考有用而且重要的一些概念。然而, 有相当多的证据表明 CAPM 没有给出风险的完整描述。^① CAPM 的投资组合建议是什么? 当有超过一种系统性风险影响资产收益率时我们应该如何改进 CAPM? 一位相信 CAPM 能够解释资产收益率的投资者将持有一个只包含无风险资产和风险资产市场组合的投资组合。如果投资者具有很高的风险容忍程度, 那么她将持有更多比例的市场组合。但是, 对于投资者持有风险资产的比例来说, 她将根据它们的市值持有一定比例的权重, 而不考虑其他任何与风险有关的因素。当然, 实际上并不是每个人都会持有相同的风险资产组合。从实际上来说, 以 CAPM 为判断标准的投资者可能持有一个货币市场基金和一个跟踪整个市场指数的投资组合。^②

在有多于一种系统性风险来源的情况下, 普通的投资者仍可能会想要持有一个包含大量风险资产的投资组合和一个无风险资产。然而, 其他投资者可能会在考虑 CAPM 没有考虑的其他风险

① 参见 Bodie, Kane 和 Marcus (2001) 关于该实证证据的介绍。

② 被动管理是与持有单个投资组合不同的一个问题。存在持有不同于 CAPM 的指数化的投资的有效市场的论述。然而, 指数基金对于该投资者来说是合理的。

之后认为持有更少的指数基金或许更好。为了给出这一结果，让我们探讨一下为什么是这样的。例如，如之前在布曼斯特（Burmeister）等人的模型中所讨论的，经济周期是系统性风险的一个来源。存在一个经济直觉来解释为什么该风险是系统性的：^①大部分投资者拥有工作，并且他们对于经济衰退非常敏感。例如，假设一位有工作的投资者面临经济衰退的风险。如果给定投资者会考虑经济衰退风险的条件，该投资者用相同的CAPM贝塔来比较2只股票，那么他将接受更低收益率的反周期性股票，并要求周期性的股票提供一个风险溢价。相反，如果该投资者具有独立的财富，而且没有丢失工作的考虑，那么他将愿意接受经济衰退的风险。

如果拥有一份工作的普通投资者会抬高反周期性股票的价格的话，那么经济衰退的风险将被定价。另外，周期性的股票将比衰退因素没有被定价时具有更低的价格。因此，正如考科伦（Cochrane）（1999a）所指出的，投资者能够“获得与市场变动不相关风险的很高的溢价”。

这种对于风险的认识对投资组合选择是有指导意义的。普通投资者会遭受到周期性风险（它是一个被定价的因素）的负面影响。（不影响普通投资者的风险不应该被定价。）拥有工作的投资者（因此，他们会获得劳动收入）想要承担更小的周期性风险，于是会要求一个周期性风险的溢价，但是没有劳动收入的投资者将愿意承担更多的周期性风险，以获得他们所不关心风险的风险溢价。因此，在其他条件相同的情况下，一位面对低于平均经济衰退风险的投资者最好的选择是承担高于平均的商业周期风险的风险暴露。

总而言之，投资者应该知道他们面临的是哪些定价风险，并对风险暴露程度予以分析。与单因素模型相比，多因素模型能够提供投资者更加丰富的信息来寻找改善投资组合选择的方法。

① 该讨论的依据是 Cochrane（1999a）和（1999b）



附 录

Appendix A Cumulative Probabilities for a Standard Normal Distribution
 $P(Z \leq x) = N(x)$ for $x \geq 0$ or $P(Z \leq z) = N(z)$ for $z \geq 0$

x or z	0	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.00	0.5000	0.5040	0.5080	0.5120	0.5160	0.5199	0.5239	0.5279	0.5319	0.5356
0.10	0.5398	0.5438	0.5478	0.5517	0.5557	0.5596	0.5636	0.5675	0.5714	0.5753
0.20	0.5793	0.5832	0.5871	0.5910	0.5948	0.5987	0.6026	0.6064	0.6103	0.6141
0.30	0.6179	0.6217	0.6255	0.6293	0.6331	0.6368	0.6406	0.6443	0.6480	0.6517
0.40	0.6554	0.6591	0.6628	0.6664	0.6700	0.6736	0.6772	0.6808	0.6844	0.6879
0.50	0.6915	0.6950	0.6985	0.7019	0.7054	0.7088	0.7123	0.7157	0.7190	0.7224
0.60	0.7257	0.7291	0.7324	0.7357	0.7389	0.7422	0.7454	0.7486	0.7517	0.7549
0.70	0.7580	0.7611	0.7642	0.7673	0.7704	0.7734	0.7764	0.7794	0.7823	0.7852
0.80	0.7881	0.7910	0.7939	0.7967	0.7995	0.8023	0.8051	0.8078	0.8106	0.8133
0.90	0.8159	0.8186	0.8212	0.8238	0.8264	0.8289	0.8315	0.8340	0.8365	0.8389
1.00	0.8413	0.8438	0.8461	0.8485	0.8508	0.8531	0.8554	0.8577	0.8599	0.8621
1.10	0.8643	0.8665	0.8686	0.8708	0.8729	0.8749	0.8770	0.8790	0.8810	0.8830
1.20	0.8849	0.8869	0.8888	0.8907	0.8925	0.8944	0.8962	0.8980	0.8997	0.9015
1.30	0.9032	0.9049	0.9066	0.9082	0.9099	0.9115	0.9131	0.9147	0.9162	0.9177
1.40	0.9192	0.9207	0.9222	0.9236	0.9251	0.9265	0.9279	0.9292	0.9306	0.9319
1.50	0.9332	0.9345	0.9357	0.9370	0.9382	0.9394	0.9406	0.9418	0.9429	0.9441
1.60	0.9452	0.9463	0.9474	0.9484	0.9495	0.9505	0.9515	0.9525	0.9535	0.9545
1.70	0.9554	0.9564	0.9573	0.9582	0.9591	0.9599	0.9608	0.9616	0.9625	0.9633
1.80	0.9641	0.9649	0.9656	0.9664	0.9671	0.9678	0.9686	0.9693	0.9699	0.9706
1.90	0.9713	0.9719	0.9726	0.9732	0.9738	0.9744	0.9750	0.9756	0.9761	0.9767
2.00	0.9772	0.9778	0.9783	0.9788	0.9793	0.9798	0.9803	0.9808	0.9812	0.9817
2.10	0.9821	0.9826	0.9830	0.9834	0.9838	0.9842	0.9846	0.9850	0.9854	0.9857
2.20	0.9861	0.9864	0.9868	0.9871	0.9875	0.9878	0.9881	0.9884	0.9887	0.9890
2.30	0.9893	0.9896	0.9898	0.9901	0.9904	0.9906	0.9909	0.9911	0.9913	0.9916
2.40	0.9918	0.9920	0.9922	0.9925	0.9927	0.9929	0.9931	0.9932	0.9934	0.9936
2.50	0.9938	0.9940	0.9941	0.9943	0.9945	0.9946	0.9948	0.9949	0.9951	0.9952
2.60	0.9953	0.9955	0.9956	0.9957	0.9959	0.9960	0.9961	0.9962	0.9963	0.9964
2.70	0.9965	0.9966	0.9967	0.9968	0.9969	0.9970	0.9971	0.9972	0.9973	0.9974
2.80	0.9974	0.9975	0.9976	0.9977	0.9977	0.9978	0.9979	0.9979	0.9980	0.9981
2.90	0.9981	0.9982	0.9982	0.9983	0.9984	0.9984	0.9985	0.9985	0.9986	0.9986
3.00	0.9987	0.9987	0.9987	0.9988	0.9988	0.9989	0.9989	0.9989	0.9990	0.9990
3.10	0.9990	0.9991	0.9991	0.9991	0.9992	0.9992	0.9992	0.9992	0.9993	0.9993
3.20	0.9993	0.9993	0.9994	0.9994	0.9994	0.9994	0.9994	0.9995	0.9995	0.9995
3.30	0.9995	0.9995	0.9995	0.9996	0.9996	0.9996	0.9996	0.9996	0.9996	0.9997
3.40	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9998
3.50	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998
3.60	0.9998	0.9998	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.70	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.80	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.90	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000
4.00	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000

For example, to find the z-value leaving 2.5 Percent of the area/probability in the upper tail, find the element 0.9750 in the body of the table. Read 1.90 at the left end of the element's row and 0.06 at the top of the element's column, to give 1.90 + 0.06 = 1.96. Table generated with Excel.

Appendix A Cumulative Probabilities for a Standard Normal Distribution
 $P(Z \leq x) = N(x)$ for $x \leq 0$ or $P(Z \leq z) = N(z)$ for $z \leq 0$

x or Z	0	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.00	0.5000	0.4960	0.4920	0.4880	0.4840	0.4801	0.4761	0.4721	0.4681	0.4641
-0.10	0.4602	0.4562	0.4522	0.4483	0.4443	0.4404	0.4364	0.4325	0.4286	0.4247
-0.20	0.4207	0.4168	0.4129	0.4090	0.4052	0.4013	0.3974	0.3936	0.3897	0.3859
-0.30	0.3821	0.3783	0.3745	0.3707	0.3669	0.3632	0.3594	0.3557	0.3520	0.3483
-0.40	0.3446	0.3409	0.3372	0.3336	0.3300	0.3264	0.3228	0.3192	0.3156	0.3121
-0.50	0.3085	0.3050	0.3015	0.2981	0.2946	0.2912	0.2877	0.2843	0.2810	0.2776
-0.60	0.2743	0.2709	0.2676	0.2643	0.2611	0.2578	0.2546	0.2514	0.2483	0.2451
-0.70	0.2420	0.2389	0.2358	0.2327	0.2296	0.2266	0.2236	0.2206	0.2177	0.2148
-0.80	0.2119	0.2090	0.2061	0.2033	0.2005	0.1977	0.1949	0.1922	0.1894	0.1867
-0.90	0.1841	0.1814	0.1788	0.1762	0.1736	0.1711	0.1685	0.1660	0.1635	0.1611
-1.00	0.1587	0.1562	0.1539	0.1515	0.1492	0.1469	0.1446	0.1423	0.1401	0.1379
-1.10	0.1357	0.1335	0.1314	0.1292	0.1271	0.1251	0.1230	0.1210	0.1190	0.1170
-1.20	0.1151	0.1131	0.1112	0.1093	0.1075	0.1056	0.1038	0.1020	0.1003	0.0985
-1.30	0.0968	0.0951	0.0934	0.0918	0.0901	0.0885	0.0869	0.0853	0.0838	0.0823
-1.40	0.0808	0.0793	0.0778	0.0764	0.0749	0.0735	0.0721	0.0708	0.0694	0.0681
-1.50	0.0668	0.0655	0.0643	0.0630	0.0618	0.0606	0.0594	0.0582	0.0571	0.0559
-1.60	0.0548	0.0537	0.0526	0.0516	0.0505	0.0495	0.0485	0.0475	0.0465	0.0455
-1.70	0.0446	0.0436	0.0427	0.0418	0.0409	0.0401	0.0392	0.0384	0.0375	0.0367
-1.80	0.0359	0.0351	0.0344	0.0336	0.0329	0.0322	0.0314	0.0307	0.0301	0.0294
-1.90	0.0287	0.0281	0.0274	0.0268	0.0262	0.0256	0.0250	0.0244	0.0239	0.0233
-2.00	0.0228	0.0222	0.0217	0.0212	0.0207	0.0202	0.0197	0.0192	0.0188	0.0183
-2.10	0.0179	0.0174	0.0170	0.0166	0.0162	0.0158	0.0154	0.0150	0.0146	0.0143
-2.20	0.0139	0.0136	0.0132	0.0129	0.0125	0.0122	0.0119	0.0116	0.0113	0.0110
-2.30	0.0107	0.0104	0.0102	0.0099	0.0096	0.0094	0.0091	0.0089	0.0087	0.0084
-2.40	0.0082	0.0080	0.0078	0.0075	0.0073	0.0071	0.0069	0.0068	0.0066	0.0064
-2.50	0.0062	0.0060	0.0059	0.0057	0.0055	0.0054	0.0052	0.0051	0.0049	0.0048
-2.60	0.0047	0.0045	0.0044	0.0043	0.0041	0.0040	0.0039	0.0038	0.0037	0.0036
-2.70	0.0035	0.0034	0.0033	0.0032	0.0031	0.0030	0.0029	0.0028	0.0027	0.0026
-2.80	0.0026	0.0025	0.0024	0.0023	0.0023	0.0022	0.0021	0.0021	0.0020	0.0019
-2.90	0.0019	0.0018	0.0018	0.0017	0.0016	0.0016	0.0015	0.0015	0.0014	0.0014
-3.00	0.0013	0.0013	0.0013	0.0012	0.0012	0.0011	0.0011	0.0011	0.0010	0.0010
-3.10	0.0010	0.0009	0.0009	0.0009	0.0008	0.0008	0.0008	0.0008	0.0007	0.0007
-3.20	0.0007	0.0007	0.0006	0.0006	0.0006	0.0006	0.0006	0.0005	0.0005	0.0005
-3.30	0.0005	0.0005	0.0005	0.0004	0.0004	0.0004	0.0004	0.0004	0.0004	0.0003
-3.40	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0002
-3.50	0.0002	0.0002	0.0002	0.0002	0.0002	0.0002	0.0002	0.0002	0.0002	0.0002
-3.60	0.0002	0.0002	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001
-3.70	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001
-3.80	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001	0.0001
-3.90	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000
-4.00	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000

For example, to find the z-value leaving 2.5 percent of the area/probability in the lower tail, find the element 0.0250 in the body of the table. Read -1.90 at the left end of the element's row and 0.06 at the top of the element's column, to give -1.90 - 0.06 = -1.96.
Table generated with Excel.

Appendix B Table of the Student's *t*-Distribution (One-Tailed Probabilities)

df	<i>p</i> = 0. 10	<i>p</i> = 0. 05	<i>p</i> = 0. 025	<i>p</i> = 0. 01	<i>p</i> = 0. 05	df	<i>p</i> = 0. 10	<i>p</i> = 0. 05	<i>p</i> = 0. 025	<i>p</i> = 0. 01	<i>p</i> = 0. 05
1	3. 078	6. 314	12. 706	31. 821	63. 657	31	1. 309	1. 696	2. 040	2. 453	2. 744
2	1. 886	2. 920	4. 303	6. 965	9. 925	32	1. 309	1. 694	2. 037	2. 449	2. 738
3	1. 638	2. 353	3. 182	4. 541	5. 841	33	1. 308	1. 692	2. 035	2. 445	2. 733
4	1. 533	2. 132	2. 776	3. 747	4. 604	34	1. 307	1. 691	2. 032	2. 441	2. 728
5	1. 476	2. 015	2. 571	3. 365	4. 032	35	1. 306	1. 690	2. 030	2. 438	2. 724
6	1. 440	1. 943	2. 447	3. 143	3. 707	36	1. 306	1. 688	2. 028	2. 434	2. 719
7	1. 415	1. 895	2. 365	2. 998	3. 499	37	1. 305	1. 687	2. 026	2. 431	2. 715
8	1. 397	1. 860	2. 306	2. 896	3. 355	38	1. 304	1. 686	2. 024	2. 429	2. 712
9	1. 383	1. 833	2. 262	2. 821	3. 250	39	1. 304	1. 685	2. 023	2. 426	2. 708
10	1. 372	1. 812	2. 228	2. 764	3. 169	40	1. 303	1. 684	2. 021	2. 423	2. 704
11	1. 363	1. 796	2. 201	2. 718	3. 106	41	1. 303	1. 683	2. 020	2. 421	2. 701
12	1. 356	1. 782	2. 179	2. 681	3. 055	42	1. 302	1. 682	2. 018	2. 418	2. 698
13	1. 350	1. 771	2. 160	2. 650	3. 012	43	1. 302	1. 681	2. 017	2. 416	2. 695
14	1. 345	1. 761	2. 145	2. 624	2. 977	44	1. 301	1. 680	2. 015	2. 414	2. 692
15	1. 341	1. 753	2. 131	2. 602	2. 947	45	1. 301	1. 679	2. 014	2. 412	2. 690
16	1. 337	1. 746	2. 120	2. 583	2. 921	46	1. 300	1. 679	2. 013	2. 410	2. 687
17	1. 333	1. 740	2. 110	2. 567	2. 898	47	1. 300	1. 678	2. 012	2. 408	2. 685
18	1. 330	1. 734	2. 101	2. 552	2. 878	48	1. 299	1. 677	2. 011	2. 407	2. 682
19	1. 328	1. 729	2. 093	2. 539	2. 861	49	1. 299	1. 677	2. 010	2. 405	2. 680
20	1. 325	1. 725	2. 086	2. 528	2. 845	50	1. 299	1. 676	2. 009	2. 403	2. 678
21	1. 323	1. 721	2. 080	2. 518	2. 831	60	1. 296	1. 671	2. 000	2. 390	2. 660
22	1. 321	1. 717	2. 074	2. 508	2. 819	70	1. 294	1. 667	1. 994	2. 381	2. 648
23	1. 319	1. 714	2. 069	2. 500	2. 807	80	1. 292	1. 664	1. 990	2. 374	2. 639
24	1. 318	1. 711	2. 064	2. 492	2. 797	90	1. 291	1. 662	1. 987	2. 368	2. 632
25	1. 316	1. 708	2. 060	2. 485	2. 787	100	1. 290	1. 660	1. 984	2. 364	2. 626
26	1. 315	1. 706	2. 056	2. 479	2. 779	110	1. 289	1. 659	1. 982	2. 361	2. 621
27	1. 314	1. 703	2. 052	2. 473	2. 771	120	1. 289	1. 658	1. 980	2. 358	2. 617
28	1. 313	1. 701	2. 048	2. 467	2. 763	200	1. 286	1. 653	1. 972	2. 345	2. 601
29	1. 311	1. 699	2. 045	2. 462	2. 756	∞	1. 282	1. 645	1. 960	2. 326	2. 576
30	1. 310	1. 697	2. 042	2. 457	2. 750						

To find a critical *t*-value, enter the with df and a specified value α , the significance level. For example, with 5 df, $\alpha = 0. 05$ and a one-tailed test, the desired probability in the tail would be $p = 0. 05$ and the critical *t*-value would be $t(5, 0. 05) = 2. 015$. With $\alpha = 0. 05$ and a two-tailed test, the desired probability in each tail would be $p = 0. 025 = \alpha/2$, giving $t(0. 025) = 2. 571$. Table generated with Excel.

Appendix C Values χ^2 (Degrees of Freedom, Level of Significance)

Degrees of Freedom	Probability in Right Tail								
	0. 99	0. 975	0. 95	0. 9	0. 1	0. 05	0. 025	0. 01	0. 005
1	0. 000 157	0. 000 982	0. 003 932	0. 015 8	2. 706	3. 841	5. 024	6. 635	7. 879
2	0. 020 100	0. 050 636	0. 102 586	0. 210 7	4. 605	5. 991	7. 378	9. 210	10. 597
3	0. 114 8	0. 215 8	0. 351 8	0. 584 4	6. 251	7. 815	9. 348	11. 345	12. 838
4	0. 297	0. 484	0. 711	1. 064	7. 779	9. 488	11. 143	13. 277	14. 860
5	0. 554	0. 831	1. 145	1. 610	9. 236	11. 070	12. 832	15. 086	16. 750
6	0. 872	1. 237	1. 635	2. 204	10. 645	12. 592	14. 449	16. 812	18. 548
7	1. 239	1. 690	2. 167	2. 833	12. 017	14. 067	16. 013	18. 475	20. 278
8	1. 647	2. 180	2. 733	3. 490	13. 362	15. 507	17. 535	20. 090	21. 955
9	2. 088	2. 700	3. 325	4. 168	14. 684	16. 919	19. 023	21. 666	23. 589
10	2. 558	3. 247	3. 940	4. 865	15. 987	18. 307	20. 483	23. 209	25. 188
11	3. 053	3. 816	4. 575	5. 578	17. 275	19. 675	21. 920	24. 725	26. 757
12	3. 571	4. 404	5. 226	6. 304	18. 549	21. 026	23. 337	26. 217	28. 300
13	4. 107	5. 009	5. 892	7. 041	19. 812	22. 362	24. 736	27. 688	29. 819
14	4. 660	5. 629	6. 571	7. 790	21. 064	23. 685	26. 119	29. 141	31. 319
15	5. 229	6. 262	7. 261	8. 547	22. 307	24. 996	27. 488	30. 578	32. 801
16	5. 812	6. 908	7. 962	9. 312	23. 542	26. 296	28. 845	32. 000	34. 267
17	6. 408	7. 564	8. 672	10. 085	24. 769	27. 587	30. 191	33. 409	35. 718
18	7. 015	8. 231	9. 390	10. 865	25. 989	28. 869	31. 526	34. 805	37. 156
19	7. 633	8. 907	10. 117	11. 651	27. 204	30. 144	32. 852	36. 191	38. 582
20	8. 260	9. 591	10. 851	12. 443	28. 412	31. 410	34. 170	37. 566	39. 997
21	8. 897	10. 283	11. 591	13. 240	29. 615	32. 671	35. 479	38. 932	41. 401
22	9. 542	10. 982	12. 338	14. 041	30. 813	33. 924	36. 781	40. 289	42. 796
23	10. 196	11. 689	13. 091	14. 848	32. 007	35. 172	38. 076	41. 638	44. 181
24	10. 856	12. 401	13. 848	15. 659	33. 196	36. 415	39. 364	42. 980	45. 558
25	11. 524	13. 120	14. 611	16. 473	34. 382	37. 652	40. 646	44. 314	46. 928
26	12. 198	13. 844	15. 379	17. 292	35. 563	38. 885	41. 923	45. 642	48. 290
27	12. 878	14. 573	16. 151	18. 114	36. 741	40. 113	43. 195	46. 963	49. 645
28	13. 565	15. 308	16. 928	18. 939	37. 916	41. 337	44. 461	48. 278	50. 994
29	14. 256	16. 047	17. 708	19. 768	39. 087	42. 557	45. 722	49. 588	52. 335
30	14. 953	16. 791	18. 493	20. 599	40. 256	43. 773	46. 979	50. 892	53. 672
50	29. 707	32. 357	34. 764	37. 689	63. 167	67. 505	71. 420	76. 154	79. 490
60	37. 485	40. 482	43. 188	46. 459	74. 397	79. 082	83. 298	88. 379	91. 952
80	53. 540	57. 153	60. 391	64. 278	96. 578	101. 879	106. 629	112. 329	116. 321
100	70. 065	74. 222	77. 929	82. 358	118. 498	124. 342	129. 561	135. 807	140. 170

To have a probability of 0. 05 in the right tail when $df = 5$, the trabled value is $\chi^2 (5, 0. 05) = 11. 070$.

Appendix D Table of the *F*-Distribution

Panel A. Critical values for right-hand tail area equal to 0.025

df1:		1	2	3	4	5	6	7	8	9	10	11	12	15	20	21	22	23	24	25	30	40	60	120	∞
df2:		1	2	3	4	5	6	7	8	9	10	11	12	15	20	21	22	23	24	25	30	40	60	120	∞
1	161	18.5	19.0	19.2	19.2	19.3	19.3	19.4	19.4	19.4	19.4	19.4	19.4	19.4	19.4	19.4	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5
2	10.1	9.55	9.28	9.12	9.12	9.01	8.94	8.89	8.85	8.81	8.79	8.76	8.74	8.70	8.66	8.65	8.65	8.64	8.64	8.63	8.62	8.59	8.57	8.55	8.53
3	7.71	6.94	6.59	6.39	6.26	6.16	6.09	6.04	6.00	5.96	5.94	5.94	5.91	5.86	5.80	5.79	5.79	5.78	5.77	5.77	5.75	5.72	5.69	5.66	5.63
4	6.61	5.79	5.41	5.19	5.05	4.95	4.88	4.82	4.77	4.74	4.74	4.70	4.68	4.62	4.56	4.55	4.54	4.53	4.53	4.52	4.50	4.46	4.43	4.40	4.37
5	5.99	5.14	4.76	4.53	4.39	4.28	4.21	4.15	4.10	4.06	4.06	4.03	4.00	3.94	3.87	3.86	3.86	3.85	3.84	3.83	3.81	3.77	3.74	3.70	3.67
6	5.59	4.74	4.35	4.12	3.97	3.87	3.79	3.73	3.68	3.64	3.64	3.60	3.57	3.51	3.44	3.43	3.43	3.42	3.41	3.40	3.38	3.34	3.30	3.27	3.23
7	5.32	4.46	4.07	3.84	3.69	3.58	3.50	3.44	3.39	3.35	3.35	3.31	3.28	3.22	3.15	3.14	3.13	3.12	3.12	3.11	3.08	3.04	3.01	2.97	2.93
8	5.12	4.26	3.86	3.63	3.48	3.37	3.29	3.23	3.18	3.14	3.14	3.10	3.07	3.01	2.94	2.93	2.92	2.91	2.90	2.89	2.86	2.83	2.79	2.75	2.71
9	4.96	4.10	3.71	3.48	3.33	3.22	3.14	3.07	3.02	2.98	2.98	2.94	2.91	2.85	2.77	2.76	2.75	2.75	2.74	2.73	2.70	2.66	2.62	2.58	2.54
10	4.84	3.98	3.59	3.36	3.20	3.09	3.01	2.95	2.90	2.85	2.85	2.82	2.79	2.72	2.65	2.64	2.63	2.62	2.61	2.60	2.57	2.53	2.49	2.45	2.40
11	4.75	3.89	3.49	3.26	3.11	3.00	2.91	2.85	2.80	2.75	2.75	2.72	2.69	2.62	2.54	2.53	2.52	2.51	2.51	2.50	2.47	2.43	2.38	2.34	2.30
12	4.67	3.81	3.41	3.18	3.03	2.92	2.83	2.77	2.71	2.67	2.67	2.63	2.60	2.53	2.46	2.45	2.44	2.43	2.42	2.41	2.38	2.34	2.30	2.25	2.21
13	4.60	3.74	3.34	3.11	2.96	2.85	2.76	2.70	2.65	2.60	2.60	2.57	2.53	2.46	2.39	2.38	2.37	2.36	2.35	2.34	2.31	2.27	2.22	2.18	2.13
14	4.54	3.68	3.29	3.06	2.90	2.79	2.71	2.64	2.59	2.54	2.54	2.51	2.48	2.40	2.33	2.32	2.31	2.30	2.29	2.28	2.25	2.20	2.16	2.11	2.07
15	4.49	3.63	3.24	3.01	2.85	2.74	2.66	2.59	2.54	2.49	2.49	2.46	2.42	2.35	2.28	2.26	2.25	2.24	2.24	2.23	2.19	2.15	2.11	2.06	2.01
16	4.45	3.59	3.20	2.96	2.81	2.70	2.61	2.55	2.49	2.45	2.45	2.41	2.38	2.31	2.23	2.22	2.21	2.20	2.19	2.18	2.15	2.10	2.06	2.01	1.96
17	4.41	3.55	3.16	2.93	2.77	2.66	2.58	2.51	2.46	2.41	2.41	2.37	2.34	2.27	2.19	2.18	2.17	2.16	2.15	2.14	2.11	2.06	2.02	1.97	1.92
18	4.38	3.52	3.13	2.90	2.74	2.63	2.54	2.48	2.42	2.38	2.38	2.34	2.31	2.23	2.16	2.14	2.13	2.12	2.11	2.11	2.07	2.03	1.98	1.93	1.88
19	4.35	3.49	3.10	2.87	2.71	2.60	2.51	2.45	2.39	2.35	2.35	2.31	2.28	2.20	2.12	2.11	2.10	2.09	2.08	2.07	2.04	1.99	1.95	1.90	1.84
20	4.32	3.47	3.07	2.84	2.68	2.57	2.49	2.42	2.37	2.32	2.32	2.28	2.25	2.18	2.10	2.08	2.07	2.06	2.05	2.05	2.01	1.96	1.92	1.87	1.81
21	4.30	3.44	3.05	2.82	2.66	2.55	2.46	2.40	2.34	2.30	2.30	2.26	2.23	2.15	2.07	2.06	2.05	2.04	2.03	2.02	1.98	1.94	1.89	1.84	1.78
22	4.28	3.42	3.03	2.80	2.64	2.53	2.44	2.37	2.32	2.27	2.27	2.24	2.20	2.13	2.05	2.04	2.02	2.01	2.01	2.00	1.96	1.91	1.86	1.81	1.76
23	4.26	3.40	3.01	2.78	2.62	2.51	2.42	2.36	2.30	2.25	2.25	2.22	2.18	2.11	2.03	2.01	2.00	1.99	1.98	1.97	1.94	1.89	1.84	1.79	1.73
24	4.24	3.39	2.99	2.76	2.60	2.49	2.40	2.34	2.28	2.24	2.24	2.20	2.16	2.09	2.01	2.00	1.98	1.97	1.96	1.96	1.92	1.87	1.82	1.77	1.71
25	4.17	3.32	2.92	2.69	2.53	2.42	2.33	2.27	2.21	2.16	2.16	2.13	2.09	2.01	1.93	1.92	1.91	1.90	1.89	1.88	1.84	1.79	1.74	1.68	1.62
30	4.08	3.23	2.84	2.61	2.45	2.34	2.25	2.18	2.12	2.08	2.08	2.04	2.00	1.92	1.84	1.83	1.81	1.80	1.79	1.78	1.74	1.69	1.64	1.58	1.51
40	4.00	3.15	2.76	2.53	2.37	2.25	2.17	2.10	2.04	1.99	1.99	1.95	1.92	1.84	1.75	1.73	1.72	1.71	1.70	1.69	1.65	1.59	1.53	1.47	1.39
60	3.92	3.07	2.68	2.45	2.29	2.18	2.09	2.02	1.96	1.91	1.91	1.87	1.83	1.75	1.66	1.64	1.63	1.62	1.61	1.60	1.55	1.50	1.43	1.35	1.25
120	3.84	3.00	2.60	2.37	2.21	2.10	2.01	1.94	1.88	1.83	1.83	1.79	1.75	1.67	1.57	1.56	1.54	1.53	1.52	1.51	1.46	1.39	1.32	1.22	1.00
Infinity																									

With 1 degree of freedom(df) in the numerator and 3 df in the denominator, the critical *F*-value is 10.1 for a right-hand tail area equal to 0.05.

Appendix D Table of the *F*-Distribution

Panel C. Critical values for right-hand tail area equal to 0. 01

Panel C. Critical values for right-hand tail area equal to 0. 01																										Numerator df ₁ and Denominator: df ₂									
df1:	1	2	3	4	5	6	7	8	9	10	11	12	15	20	21	22	23	24	25	30	40	60	120	∞											
df2:1	405.2	500.0	540.3	562.5	576.4	585.9	592.8	598.2	602.3	605.6	608.3	610.6	615.7	620.9	621.6	622.3	622.9	623.5	624.0	626.1	628.7	631.3	633.9	636.6											
2	98.5	99.0	99.2	99.2	99.3	99.3	99.4	99.4	99.4	99.4	99.4	99.4	99.4	99.4	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5											
3	34.1	30.8	29.5	28.7	28.2	27.9	27.7	27.5	27.3	27.2	27.1	27.1	26.9	26.7	26.7	26.6	26.6	26.6	26.6	26.5	26.4	26.3	26.2	26.1											
4	21.2	18.0	16.7	16.0	15.5	15.2	15.0	14.8	14.7	14.5	14.5	14.4	14.2	14.0	14.0	14.0	13.9	13.9	13.9	13.8	13.7	13.7	13.6	13.5											
5	16.3	13.3	12.1	11.4	11.0	10.7	10.5	10.3	10.2	10.1	10.0	9.89	9.72	9.55	9.53	9.51	9.49	9.47	9.45	9.38	9.29	9.20	9.11	9.02											
6	13.7	10.9	9.78	9.15	8.75	8.47	8.26	8.10	7.98	7.87	7.79	7.72	7.56	7.40	7.37	7.35	7.33	7.31	7.30	7.23	7.14	7.06	6.97	6.88											
7	12.2	9.55	8.45	7.85	7.46	7.19	6.99	6.84	6.72	6.62	6.54	6.47	6.31	6.16	6.13	6.11	6.09	6.07	6.06	5.99	5.91	5.82	5.74	5.65											
8	11.3	8.65	7.59	7.01	6.63	6.37	6.18	6.03	5.91	5.81	5.73	5.67	5.52	5.36	5.34	5.32	5.30	5.28	5.26	5.20	5.12	5.03	4.95	4.86											
9	10.6	8.02	6.99	6.42	6.06	5.80	5.61	5.47	5.35	5.26	5.18	5.11	4.96	4.81	4.79	4.77	4.75	4.73	4.71	4.65	4.57	4.48	4.40	4.31											
10	10.0	7.56	6.55	5.99	5.64	5.39	5.20	5.06	4.94	4.85	4.77	4.71	4.56	4.41	4.38	4.36	4.34	4.33	4.31	4.25	4.17	4.08	4.00	3.91											
11	9.65	7.21	6.22	5.67	5.32	5.07	4.89	4.74	4.63	4.54	4.46	4.40	4.25	4.10	4.08	4.06	4.04	4.02	4.01	3.94	3.86	3.78	3.69	3.60											
12	9.33	6.93	5.95	5.41	5.06	4.82	4.64	4.50	4.39	4.30	4.22	4.16	4.01	3.86	3.84	3.82	3.80	3.78	3.76	3.70	3.62	3.54	3.45	3.36											
13	9.07	6.70	5.74	5.21	4.86	4.62	4.44	4.30	4.19	4.10	4.02	3.96	3.82	3.66	3.64	3.62	3.60	3.59	3.57	3.51	3.43	3.34	3.25	3.17											
14	8.86	6.51	5.56	5.04	4.70	4.46	4.28	4.14	4.03	3.94	3.86	3.80	3.66	3.51	3.48	3.46	3.44	3.43	3.41	3.35	3.27	3.18	3.09	3.00											
15	8.68	6.36	5.42	4.89	4.56	4.32	4.14	4.00	3.89	3.80	3.73	3.67	3.52	3.37	3.35	3.33	3.31	3.29	3.28	3.21	3.13	3.05	2.96	2.87											
16	8.53	6.23	5.29	4.77	4.44	4.20	4.03	3.89	3.78	3.69	3.62	3.55	3.41	3.26	3.24	3.22	3.20	3.18	3.16	3.10	3.02	2.93	2.84	2.75											
17	8.40	6.11	5.19	4.67	4.34	4.10	3.93	3.79	3.68	3.59	3.52	3.46	3.31	3.16	3.14	3.12	3.10	3.08	3.07	3.00	2.92	2.83	2.75	2.65											
18	8.29	6.01	5.09	4.58	4.25	4.01	3.84	3.71	3.60	3.51	3.43	3.37	3.23	3.08	3.05	3.03	3.02	3.00	2.98	2.92	2.84	2.75	2.66	2.57											
19	8.19	5.93	5.01	4.50	4.17	3.94	3.77	3.63	3.52	3.43	3.36	3.30	3.15	3.00	2.98	2.96	2.94	2.92	2.91	2.84	2.76	2.67	2.58	2.49											
20	8.10	5.85	4.94	4.43	4.10	3.87	3.70	3.56	3.46	3.37	3.29	3.23	3.09	2.94	2.92	2.90	2.88	2.86	2.84	2.78	2.69	2.61	2.52	2.42											
21	8.02	5.78	4.87	4.37	4.04	3.81	3.64	3.51	3.40	3.31	3.24	3.17	3.03	2.88	2.86	2.84	2.82	2.80	2.79	2.72	2.64	2.55	2.46	2.36											
22	7.95	5.72	4.82	4.31	3.99	3.76	3.59	3.45	3.35	3.26	3.18	3.12	2.98	2.83	2.81	2.78	2.77	2.75	2.73	2.67	2.58	2.50	2.40	2.31											
23	7.88	5.66	4.76	4.26	3.94	3.71	3.54	3.41	3.30	3.21	3.14	3.07	2.93	2.78	2.76	2.74	2.72	2.70	2.69	2.62	2.54	2.45	2.35	2.26											
24	7.82	5.61	4.72	4.22	3.90	3.67	3.50	3.36	3.26	3.17	3.09	3.03	2.89	2.74	2.72	2.70	2.68	2.66	2.64	2.58	2.49	2.40	2.31	2.21											
25	7.77	5.57	4.68	4.18	3.86	3.63	3.46	3.32	3.22	3.13	3.06	2.99	2.85	2.70	2.68	2.66	2.64	2.62	2.60	2.53	2.45	2.36	2.27	2.17											
30	7.56	5.39	4.51	4.02	3.70	3.47	3.30	3.17	3.07	2.98	2.91	2.84	2.70	2.55	2.53	2.51	2.49	2.47	2.45	2.39	2.30	2.21	2.11	2.01											
40	7.31	5.18	4.31	3.83	3.51	3.29	3.12	2.99	2.89	2.80	2.73	2.66	2.52	2.37	2.35	2.33	2.31	2.29	2.27	2.20	2.11	2.02	1.92	1.80											
60	7.08	4.98	4.13	3.65	3.34	3.12	2.95	2.82	2.72	2.63	2.56	2.50	2.35	2.20	2.17	2.15	2.13	2.12	2.10	2.03	1.94	1.84	1.73	1.60											
120	6.85	4.79	3.95	3.48	3.17	2.96	2.79	2.66	2.56	2.47	2.40	2.34	2.19	2.03	2.01	1.99	1.97	1.95	1.93	1.86	1.76	1.66	1.53	1.38											
Infinity	6.63	4.61	3.78	3.32	3.02	2.80	2.64	2.51	2.41	2.32	2.25	2.18	2.04	1.88	1.85	1.83	1.81	1.79	1.77	1.70	1.59	1.47	1.32	1.00											

Appendix E Critical Values for the Durbin-Watson Statistic ($\alpha=0.05$)

<i>n</i>	<i>K</i> = 1		<i>K</i> = 2		<i>K</i> = 3		<i>K</i> = 4		<i>K</i> = 5	
	<i>d_L</i>	<i>d_U</i>	<i>d_L</i>	<i>d_U</i>	<i>d_L</i>	<i>d_U</i>	<i>d_L</i>	<i>d_U</i>	<i>d_L</i>	<i>d_U</i>
15	1.08	1.36	0.95	1.54	0.82	1.75	0.69	1.97	0.56	2.21
16	1.10	1.37	0.98	1.54	0.86	1.73	0.74	1.93	0.62	2.15
17	1.13	1.38	1.02	1.54	0.90	1.71	0.78	1.90	0.67	2.10
18	1.16	1.39	1.05	1.53	0.93	1.69	0.82	1.87	0.71	2.06
19	1.18	1.40	1.08	1.53	0.97	1.68	0.86	1.85	0.75	2.02
20	1.20	1.41	1.10	1.54	1.00	1.68	0.90	1.83	0.79	1.99
21	1.22	1.42	1.13	1.54	1.03	1.67	0.93	1.81	0.83	1.96
22	1.24	1.43	1.15	1.54	1.05	1.66	0.96	1.80	0.86	1.94
23	1.26	1.44	1.17	1.54	1.08	1.66	0.99	1.79	0.90	1.92
24	1.27	1.45	1.19	1.55	1.10	1.66	1.01	1.78	0.93	1.90
25	1.29	1.45	1.21	1.55	1.12	1.66	1.04	1.77	0.95	1.89
26	1.30	1.46	1.22	1.55	1.14	1.65	1.06	1.76	0.98	1.88
27	1.32	1.47	1.24	1.56	1.16	1.65	1.08	1.76	1.01	1.86
28	1.33	1.48	1.26	1.56	1.18	1.65	1.10	1.75	1.03	1.85
29	1.34	1.48	1.27	1.56	1.20	1.65	1.12	1.74	1.05	1.84
30	1.35	1.49	1.28	1.57	1.21	1.65	1.14	1.74	1.07	1.83
31	1.36	1.50	1.30	1.57	1.23	1.65	1.16	1.74	1.09	1.83
32	1.37	1.50	1.31	1.57	1.24	1.65	1.18	1.73	1.11	1.82
33	1.38	1.51	1.32	1.58	1.26	1.65	1.19	1.73	1.13	1.81
34	1.39	1.51	1.33	1.58	1.27	1.65	1.21	1.73	1.15	1.81
35	1.40	1.52	1.34	1.58	1.28	1.65	1.22	1.73	1.16	1.80
36	1.41	1.52	1.35	1.59	1.29	1.65	1.24	1.73	1.18	1.80
37	1.42	1.53	1.36	1.59	1.31	1.66	1.25	1.72	1.19	1.80
38	1.43	1.54	1.37	1.59	1.32	1.66	1.26	1.72	1.21	1.79
39	1.43	1.54	1.38	1.60	1.33	1.66	1.27	1.72	1.22	1.79
40	1.44	1.54	1.39	1.60	1.34	1.66	1.29	1.72	1.23	1.79
45	1.48	1.57	1.43	1.62	1.38	1.67	1.34	1.72	1.29	1.78
50	1.50	1.59	1.46	1.63	1.42	1.67	1.38	1.72	1.34	1.77
55	1.53	1.60	1.49	1.64	1.45	1.68	1.41	1.72	1.38	1.77
60	1.55	1.62	1.51	1.65	1.48	1.69	1.44	1.73	1.41	1.77
65	1.57	1.63	1.54	1.66	1.50	1.70	1.47	1.73	1.44	1.77
70	1.58	1.64	1.55	1.67	1.52	1.70	1.49	1.74	1.46	1.77
75	1.60	1.65	1.57	1.68	1.54	1.71	1.51	1.74	1.49	1.77
80	1.61	1.66	1.59	1.69	1.56	1.72	1.53	1.74	1.51	1.77
85	1.62	1.67	1.60	1.70	1.57	1.72	1.55	1.75	1.52	1.77
90	1.63	1.68	1.61	1.70	1.59	1.73	1.57	1.75	1.54	1.78
95	1.64	1.69	1.62	1.71	1.60	1.73	1.58	1.75	1.56	1.78
100	1.65	1.69	1.63	1.72	1.61	1.74	1.59	1.76	1.57	1.78

Soure: From J. Durbin and G. S. Watson, "Testing for Serial Correlation in Least Squares Regressinon, II." *Biometrika* 38 (1951): 159—178. Reproduced by permission of the *Biometrika* trustees.

Note: *K* = the number of slope parameters in the model.

术 语 表



a priori probability **先验概率** 基于逻辑分析而非观测所形成的概率或者是由个人主观判断而得到的概率。

absolute dispersion **绝对离散度** 没有与任何参考点或基准进行比较而得到的反映数据变动程度的量。

absolute frequency **绝对频数** 对数据集进行分组后,落在某个给定区间内的数据个数。

accrued interest **应计利息** (会计上)已视为收入但还未实际获得的利息。

active factor risk **积极因素风险** 由相对于风险模型中具体因素的投资组合不同于基准的风险暴露所带来的对于积极风险平方的贡献。

active return **积极收益率** 投资组合收益率超出投资组合基准收益率的部分(投资组合的收益率减去投资组合基准的收益率)。

active risk **积极风险** 积极收益率的标准差。

active risk squared **积极风险平方** 积极收益率的方差;积极风险的二次方。

active specific risk/asset selection risk **积极特有风险/资产选择风险** 由于投资组合对各个资产赋予积极权重而产生的(因为这些权重与资产的残差风险互相影响)对于积极风险平方的贡献。

addition rule for probabilities **概率的加法法则**
计算事件 A 和 B 中至少有一个发生(包括同时发生)的概率法则,即事件 A 和 B 中至少有一个发生的概率等于事件 A 发生的概率加上事件 B 发生的概率减去事件 A 和 B 同时发生的概率。

adjusted beta **调整后的贝塔** 为反映贝塔均值回复的趋势而对其历史值进行调整后所获得的贝塔。

adjusted R^2 **调整后的 R 平方** 反映回归拟合度的一种度量。由于其对自由度进行了调整,因此,当回归中的自变量增加时,其值并不会自动增加。

alternative hypothesis **备择假设** 当原假设被拒绝时所接受的假设。

analysis of variance, ANOVA **方差分析** 将数据集(如回归中因变量的观测值)的总体变动分解为代表不同变动来源的组成成分的一种分析方法;在回归分析中,方差分析提供了回归总体显著性 F 检验中所使用的输入变量。

annual percentage rate **年度百分比利率** 以年利率表示的借款成本。

annuity **年金** 一组由有限个等额、等时间间隔的现金流所组成的序列。

annuity due **先付年金** 首笔现金流(在期初)即刻支付的年金。

arbitrage **套利** 一种无风险的交易策略,它在不需资金净投资的情况下就能获得正的期望净收益。

arbitrage opportunity **套利机会** 进行套利的机会;在不承担风险、没有资金净投资的情况下,获取正的期望净收益的机会。

arbitrage portfolio **套利(投资)组合** 能够获得套利机会的投资组合。

arithmetic mean **算术平均数** 所有观测值的总和除以观测值的数目之后所得到的数值。

asian call option **亚式看涨期权** 欧式期权的一种,这种期权在到期日的价值等于股票在到期日的价格和期权有效期(生存期)内股票平均价格的差额与0之间的较大者。

autocorrelation **自相关系数** 时间序列中的随机变量与其自身历史值之间的相关系数。

autoregressive model, AR **自回归模型** 用其自身的历史值进行回归的时间序列模型, 它的自变量是因变量的滞后值。

bank discount basis **银行贴现基准** 一种报价的惯例, 它将贴现率按每年 360 天来年化为面值的某个百分比。

bayes' formula **贝叶斯公式** 基于新信息对概率进行更新的一种方法。

benchmark **基准** 一个用来作为参照的投资组合; 用来作为参考或者比较的一个点。

bernoulli random variable **伯努利随机变量** 只具有 0 和 1 两种结果的随机变量。

bernoulli trial **伯努利试验** 只会产生两种结果之中一种的试验。

beta **贝塔** 一个度量资产对于市场变化敏感程度的量。

binomial random variable **二项随机变量** n 次伯努利试验中成功的次数, 其中所有伯努利试验成功的概率为常数, 并且各次试验之间是相互独立的。

binomial tree **二叉树** 资产价格动态变化模型的一种图形表示, 在每个节点上, 资产价格以概率 p 上升或以概率 $1 - p$ 下降。

block **大单** 超出做市商网络或者股票交易所所能提供的正常流动性的买单或卖单。

bond-equivalent basis **债券等价基准** 一种通过将半年收益率翻倍进行年化的年收益率报价的基准。

bond-equivalent yield **债券等值收益** 不考虑复利情况下的到期收益率。

breusch-Pagan test **布鲁奇—帕甘检验** 一种检验回归中误差项条件异方差的方法。

capital allocation line, CAL **资本配置线** 在图形上反映投资者最优投资组合中风险资产和无风险资产各种可行的期望收益率和收益率标准差组合全体的一条直线。

capital asset pricing model, CAPM **资本资产定价模型** 将任何资产 (或投资组合) 的期望收益率表示为关于其贝塔的线性函数的一个等式。

capital budgeting **资本预算** 将资金分配到相对长期的项目或者投资中的行为。

capital market line, CML **资本市场线** 资本配置线的一种形式, 它假设了投资者对于风险资产收益率的均值、方差以及相关性具有相同的预期。

capital structure **资本结构** 一个公司长期融资方式的具体构成。

cash flow additivity principle **现金流加法原理** 在同一时点上的美元数量具有可加性的原理。

central limit theorem **中心极限定理** 统计学的一项基本定理。根据此定理, 来自有限方差总体的样本量为 n 的大样本所计算出的样本均值, 将会近似服从一个均值为总体均值、方差为总体方差除以 n 的正态分布。

chain rule of forecasting **预测的链式法则** 将下一期的预测值代入预测方程的右边来预测两期之后数值的预测过程。

coefficient of variation, CV **变异系数** 一组观测值的标准差与其均值的比值。

cointegrated **协整** 用来描述两个具有长期金融或经济关系并且在长期中相互不会偏离很大的时间序列。

combination **组合** 数据的一个排序, 其中数据项的先后次序无关紧要。

commercial paper **商业票据** 无担保的短期公司债务, 它以到期日一次性偿还贷款为其特征。

common size statements **同比财务报表** 财务报表中的所有元素 (账目) 都表示为某个关键数字 (如利润表中的营业收入或者是资产负债表中的总资产) 的百分比。

company fundamental factors **公司基本面因素** 那些反映公司内部业绩表现的因素, 如与盈利增长、盈利变动、盈利惯性以及财务杠杆等相关的因素。

company share-related factors **公司股价关联因素** 估值变量以及其他与公司股票价格或者股票的交易特征相关的因素, 例如收益

率、股息生息率以及账面市值比。

complement 补集 关于事件 S , 指 S 不发生的事件。

compounding 复利 利息产生利息的不断累积的过程。

conditional expected value 条件期望值 给定另一个事件发生的条件下, 所要讨论事件(变量)的期望值。

conditional heteroskedasticity 条件异方差 误差项方差的异方差性与回归方程中自变量的值有相关性。

conditional probability 条件概率 给定另一事件发生的条件下, 某一事件发生的概率。

conditional variances 条件方差 给定另一个变量结果的条件下, 某一变量的方差。

confidence interval 置信区间 是指一个区间, 该区间会在给定概率水平下包含总体所要估计的参数。

consistency 相合性 估计量的一个优良性质; 相合估计量是指随着样本量的增大, 该估计量接近总体参数真值的概率也会增大, 并趋近于 1。

consistent 相合的 关于估计量的一个概念, 它所描述的是随着样本量的增大, 估计量接近总体参数真值的概率也不断增大。

continuous random variable 连续型随机变量 可能出现结果的取值范围为实直线(所有在负无穷到正无穷之间的实数)或者是实直线的某一子集的随机变量。

continuously compounded return 连续复利收益率 1 加上持有期收益率的自然对数, 或者等价地, 等于期末价格除以期初价格(相对价格)的自然对数。

correlation 相关系数 一个在 $-1 \sim 1$ 之间的数值, 其被用来衡量两个随机变量之间联合变动程度(线性关联性)。

correlation analysis 相关性分析 对两组数据序列之间线性关系的强弱进行分析。

cost averaging 成本平均 固定数量资金的周期性投资。

covariance 协方差 两个随机变量之间联合变动程度(线性关联性)的度量指标。

covariance matrix 协方差矩阵 由协方差组成的矩阵或者方阵; 也被称为方差-协方差矩阵。

covariance stationary 协方差平稳 当时间序列在所有期间的期望值和方差都是有限常数, 并且该时间序列过去或将来的任一固定期间数的自协方差在所有期间也都是有限常数时, 就称该时间序列是协方差平稳的。

cross-sectional data 截面数据 在同一时间点上对于多个个体的观测数据, 与时间序列数据相对。

cumulative distribution function 累积分布函数 给出一个随机变量小于或等于某给定数值概率的函数。

cumulative relative frequency 累积相对频数(累积频率) 对于划分为区间的数据, 小于给定区间上限的观测值占总观测值的比例。

data mining 数据挖掘 通过对数据集的广泛研究中所寻找出的统计意义上显著的模式来确定模型的一种行为。

deciles 十分位数 把分布分成十等分的分位数。

decision rule 决策规则 与假设检验有关的概念, 对是否拒绝原假设做出判断的规则; 它涉及将检验统计量与拒绝原假设的临界点进行比较。

default risk premium 违约风险溢价 对于借款人在合约到期时存在不能及时或全额还款的可能性, 投资者所要求给予的额外收益补偿。

degree of confidence 置信度 置信区间包含未知总体参数的概率。

degrees of freedom, df 自由度 使用的独立观测值的个数。

dependent 不独立/依赖 与事件有关的概念, 一个事件发生的概率依赖于另一个事件发生的概率(或与其有关)的性质。

dependent variable 因变量 该变量偏离均值的变动可以由回归来解释；它位于回归方程等号的左侧。

descriptive statistics 描述性统计 研究如何对数据进行有效的概括。

diffuse prior 原始（缺乏信息时的）先验概率假设 等先验概率假设。

discount 折价/贴现/贴水 根据未来支付在时间上的远近，来降低其价值作为补偿；计算未来某笔现金流的现值；也指金融工具的标价低于其面值的情况。

discrete random variable 离散型随机变量 一个至多能取可数个可能值的随机变量。

discriminant analysis 判别分析 一种用来区别不同组别的多变量分类技术，如划分公司是否会在某个期限内破产。

dispersion 离散度 反映数据相对于中心趋势的波动情况。

down transition probability 向下转变概率 在资产价格动态变化模型中，资产价值下降的概率。

dummy variable 虚拟变量 一种定性变量，当特定的条件成立时，它取值为1，而当条件不成立时，它取值为0。

dutch Book Theorem 荷兰赌定理 概率论中关于由不相容概率而产生获利机会的结果。

effective annual rate, EAR 有效（实际）年利率 一单位货币在一年中的增长量，包括利息所产生的利息部分。

effective annual yield, EAY 有效年收益率 考虑复利影响的年化收益率；有效年收益率等于1加上持有期收益率并复利至一年，再减去1。

efficiency 有效性 估计量的一个良好性质。一个有效估计量是指同一参数所有的无偏估计量中方差最小的那个无偏估计量。

efficient frontier 有效前沿/有效边界 最小方差前沿从全局最小方差组合开始向上延伸的部分；在由任意给定方差水平下，拥有最大期望收益率的投资组合所构成的边界图像。

efficient portfolio 有效投资组合 给定以收益率的方差或标准差度量的风险水平下，使期望收益率最大的投资组合。

empirical probability 经验概率 根据事件发生的相对频数而估计出的事件发生的概率。

error autocorrelation 误差项自相关系数 回归模型中误差项的自相关系数。

error term 误差项 因变量不能由回归方程中自变量所解释的部分。

estimate 估计值 将样本观测值代入估计式所计算出的特定值。

estimated (or fitted) parameters 估计（或拟合）参数 关于回归分析的概念，指的是回归方程中所估计出的总体截距与总体斜率系数。

estimation 估计 关于统计推断的概念，其关注的是如何估计总体参数的数值。

estimator 估计式 一个估计公式；该公式被用来计算诸如样本均值和其他样本统计量。

european-style option or European option 欧式期权或者欧洲期权 只能在到期日行使权利的期权。

event 事件 一个随机变量的任一结果或者某些特定结果的集合。

excess kurtosis 超额峰度 尖峰（厚尾）超出正态分布尖峰的程度。

exhaustive 穷举法 列举或包括所有可能的试验结果。

expected value 期望值 对随机变量所有可能取值的概率加权平均值。

face value 面值 与任何息票支付无关的、在到期日承诺支付的金额。

factor 因素 一个共同的或基础的与多个变量相关的元素。

factor risk premium (or factor price) 因素风险溢价（或因素价格） 对某一因素的敏感度为1，而同时对于其他因素的敏感度为0的投资组合超出无风险利率部分的期望收益率。

factor sensitivity (also factor betas or factor

- loadings) **因素敏感性** (又称**因素贝塔**或**因子载荷**) 在保持所有其他因素不变情况下, 某个因素增加一单位引起收益率变化程度的度量指标。
- financial risk **金融风险** 与资产价格及其他金融变量有关的风险。
- first-differencing **一阶差分** 将时间序列中 t 期的值减去 $t-1$ 期的值的一种变换。
- first-order serial correlation **一阶序列相关** 时间序列中相邻观测值之间的相关性。
- frequency distribution **频数分布** 一种表格列示数据的方法, 它用较少的区间对于数据总体进行概括。
- frequency polygon **频数多边形** 通过把代表区间频数的点用直线连接起来后, 所得到的一个表示频数分布的图形。
- full price **全价** 包含应计利息的证券价格。
- fundamental beta **基本面贝塔** 至少是部分基于公司基本面数据而得到的贝塔。
- fundamental factor models **基本面因素模型** 一种多因素模型, 在此类模型中, 因素是一些对于解释股价横截面差异有重要作用的股票或公司的属性。
- future value, FV **将来值** 到未来某个指定时刻, 一笔支付或者一系列支付将会增长到的数额。
- generalized least squares **广义最小二乘法** 一种处理误差项异方差性的回归估计技术。
- geometric mean **几何平均** 由 n 个非负数连乘积的 n 次方根计算得到的一种度量数据集中趋势的指标。
- harmonic mean **调和平均** 一类加权平均的方法, 它的计算方法为: 先对观测值的倒数计算平均, 再将该平均数求倒数。
- heteroskedastic **异方差** 关于回归误差项的概念, 误差项的方差对各个观测值是不同的。
- heteroskedasticity **异方差性** 具有非常数方差的性质; 具有各个观测值的误差项方差是不同的性质。
- heteroskedasticity-consistent standard errors **异方差一致标准误** 对回归中误差项存在异方差情况进行校正后所得到的回归估计参数的标准误。
- histogram **直方图** 用来描述由数据分类而得到的频数分布的条状图。
- historical simulation (or back simulation) **历史模拟 (或后向模拟)** 从历史数据序列中进行抽样的模拟方法。
- holding period return **持有期收益** 投资者在给定的持有期内获得的收益; 同总收益。
- holding period yield, HPY **持有期收益** 投资者在给定的持有期内获得的收益; 关于固定收益工具的持有期收益。
- homoscedasticity **同方差性** 具有常数方差的性质; 具有各个观测值的误差项方差是相同的性质。
- hurdle rate **门槛比率** 接受投资项目所必须要满足的收益率。
- hypothesis **假设** 关于统计推断的概念, 指的是关于一个或多个总体的论断。
- hypothesis testing **假设检验** 关于统计推断的概念, 涉及如何对关于一个或多个总体的假设进行检验。
- incremental cash flows **增量现金流** 由决策或投资行为引起的现金流的改变量或者增加量。
- independent **独立** 与事件有关的概念, 一个事件的发生不会影响另一个事件发生的概率的性质。
- independent and identically distributed, IID **独立同分布** 与随机变量有关的概念, 随机变量彼此独立、但服从相同的概率分布的性质。
- independent variable **自变量** 回归中用来解释因变量的变量; 它出现在回归方程等式的右边。
- indexing **指数化** 一种投资策略, 投资者构造一个能够复制某个给定指数业绩的投资组合。
- inflation premium **通货膨胀溢价** 给予投资者

的额外收益用以补偿预期通货膨胀所带来的影响。

information ratio, IR **信息比率** 积极收益率的均值除以积极风险。

in-sample forecast errors **样本内预测误差** 在拟合模型的样本期内,由拟合时间序列模型而得到的残差。

instability in the minimum-variance frontier **最小方差前沿的不稳定性** 最小方差前沿对于输入变量的小变动敏感的特征。

interest rate **利率** 反映不同时点现金流之间关系的收益率;贴现率。

intergenerational data mining **代际数据挖掘** 一种数据挖掘的形式,它应用先前研究者使用的某个数据集所获得的信息来指导当前使用相同的或相关的数据集来进行的研究。

internal rate of return, IRR **内部收益率** 使净现值等于0的贴现率;使投资成本(现金流出)的现值与投资收益(现金流入)的现值相等的贴现率。

interquartile range **四分位距** 数据集第三个四分位数与第一个四分位数之差。

interval **区间** 有关分组数据的概念,观测值落入其中的一组数值。

interval scale **间隔尺度** 一种测量尺度,它不仅对数据排序,并且保证尺度数值之间的差是相同的。

IRR rule **内部收益率准则** 内部收益率必须要大于资本机会成本才接受项目或投资方案的投资决策准则。

joint probability **联合概率** 所讨论的事件共同发生的概率。

joint probability function **联合概率函数** 给出所讨论的随机变量的值共同出现的概率的函数。

kth Order autocorrelation **K阶自相关** 时间序列中相差 k 期的两个观测值之间的相关关系。

kurtosis **峰度** 用来反映分布尖峰程度的统计度量指标。

leptokurtic **尖峰** 用来描述一个比正态分布的尖峰程度更高的分布。

level of significance **显著性水平** 假设检验中犯第一类错误的概率。

likelihood **似然** 给定一组特定条件下,某个观测值出现的概率。

linear association **线性相关** 直线关系,与不能以直线图形来反映的关系相对。

linear interpolation **线性插值** 基于连接两个已知量的直线来估计在这两个已知量之间的一个未知量的值的方法。

linear regression **线性回归** 对因变量与自变量之间的线性关系进行建模的回归。

linear trend **线性趋势** 因变量随着时间以一个常数比率变化的趋势。

liquidity premium **流动性溢价** 给予投资者的超额收益以补偿投资需要迅速转变为现金时所存在的相对于投资公允价值而言的损失的风险。

logit model **Logit 模型** 因变量是定性变量的基于 Logistic 概率分布的多元回归模型。

log-linear model **对数线性模型** 有关时间序列模型的概念,模型中时间序列的增长率作为时间的函数是常数。

log-log regression model **双对数回归模型** 因变量和自变量均表示为自然对数形式的回归模型。

longitudinal data **纵向数据** 对同一个观测单位的特性在不同时点上的观测值。

look-ahead bias **前视偏差** 由于使用了在进行检验时尚未能获得的信息而产生的偏差。

macroeconomic factor **宏观经济因素** 与宏观经济相关的因素。例如,通货膨胀率、工业产值、或行业归属指标。

macroeconomic factor model **宏观经济因素模型** 一个多因素模型,模型中因素是宏观经济变量中能够显著解释资产收益的部分。

market price of risk **风险的市场价格** 资本市场线的斜率,反映承担每一单位市场风险所要求的市场风险溢价。

market risk premium **市场风险溢价** 市场预期收益率高于无风险利率的部分。

markowitz decision rule **马科维茨决策规则** 根据两个投资项目的均值与方差来进行选择的决策规则。

maturity premium **到期日溢价** 随着到期期限的增加,债务工具的市场价值对于市场利率变化的敏感性也增加,为补偿这种风险而给予投资者的额外收益。

mean **均值** 分布或数据集中所有数值加总,再除以被加总数值的数目所得到的值;同算术平均。

mean absolute deviation **平均绝对离差** 关于样本的概念,指的是观测值对于样本均值的离差的绝对值的平均值。

mean excess return **平均超额收益** 超过无风险利率的平均收益率。

mean reversion **均值回复** 当时间序列变量水平高于其均值时有下降的趋势,而当变量水平低于其均值时有上升的趋势;具有均值回复特性的时间序列倾向于回归到其长期平均水平。

mean-variance analysis **均值方差分析** 一种利用资产收益率的期望均值、方差和协方差来对投资组合进行分析的方法。

measure of central tendency **集中趋势测度** 用来表示数据集中中心的数量测度。

measure of location **位置测度** 用来描述数据的位置或者分布的数量测度;不仅包括集中趋势测度,还包括其他测度如百分位数。

measurement scales **度量尺度** 度量差异的方法;有四种类型的度量尺度:名义尺度、顺序尺度、间隔尺度和比率尺度。

median **中位数** 将数据集升序或者降序排列后,处于中心位置的数即为中位数;第 50 个百分位数。

mesokurtic **中峰** 用来描述与正态分布的峰度相同的分布。

minimum-variance frontier **最小方差前沿** 在给定期望收益水平下具有最小方差的投资组

合集合的图形表示。

minimum-variance portfolio **最小方差投资组合** 在每个给定期望收益水平下具有最小方差的投资组合。

mixed factor models **混合因素模型** 融合了多种(一种以上)类型的因素模型的特征的因素模型。

modal interval **众数区间** 关于分组数据的概念,指发生频率最高的区间。

mode **众数** 一组观测值中出现频数最高的数值。

model specification **模型设定** 关于回归的概念,指确定回归中包含的变量以及回归方程的函数形式。

money market **货币市场** 短期(到期期限等于或少于一年)债务工具的市场。

money market yield (or CD equivalent yield) **货币市场收益(或 CD 等价收益)** 与按每年 360 天计息的有息货币市场工具的报价收益率相对应的收益率;按一年 360 天计息的年化持有期收益率。

money-weighted rate of return **货币加权收益率** 考虑所有现金流的投资组合的内部收益率。

monte Carlo simulation **蒙特卡罗模拟** 运用计算机来对复杂问题寻找近似解的一种方法。

multicollinearity **多重共线性** 一种违背回归假设的情况,它在两个或两个以上自变量(或者是自变量的组合)互相高度相关但不是完全相关时产生。

multiple linear regression **多元线性回归** 涉及两个或两个以上自变量的线性回归。

multiple linear regression model **多元线性回归模型** 拥有两个或两个以上自变量的线性回归模型。

multiple R **复相关系数** 回归中因变量的实际值与预测值之间的相关系数。

multiplication rule for probabilities **概率乘法法则** A 事件和 B 事件同时发生的概率等于已

知 B 事件发生条件下 A 事件发生的概率乘以 B 事件发生的无条件概率。

multivariate distribution 多元分布 表示一组相关随机变量概率的概率分布。

multivariate normal distribution 多元正态分布 完全由变量的均值和方差以及两两变量间相关系数决定的一组随机变量的概率分布。

mutually exclusive events 互斥事件 在某个时间点只有一个事件能发生的若干事件。

n factorial n 阶乘 对于一个正整数 n , 前 n 个正整数的乘积; 根据定义, 0 阶乘等于 1。 n 阶乘记为 $n!$

negative serial correlation 负序列相关 一种序列相关关系, 在该关系中一个观测值出现正误差会增加另一个观测值出现负误差的可能性, 反之亦然。

net present value, NPV 净现值 投资现金流入 (收入) 的现值减去投资现金流出 (成本) 的现值。

node 结点 二叉树上的每个值, 后继的移动或结果由该值处分叉形成。

nominal risk-free interest rate 名义无风险利率 实际无风险利率和通货膨胀溢价之和。

nominal scale 名义尺度 一种对数据进行分类, 但不进行排序的度量尺度。

nonlinear relation 非线性关系 不能用直线画出的变量间的联系或关系。

nonparametric test 非参数检验 一种不考虑参数或者对于样本来源的总体做出最少假设的检验方法。

nonstationarity 非平稳性 关于随机变量的概念, 均值和方差等数字特征随时间变化的性质。

normal distribution 正态分布 一种连续、对称的概率分布, 它完全由均值和方差所描述。

n -period moving average n 阶移动平均 一个时间序列的当前值和之前相邻的 $n-1$ 个值的平均值。

nPV rule 净现值准则 当净现值为正时进行投

资, 而在净现值为负时不进行投资的投资决策准则。

null hypothesis 原假设 要被检验的假设。

objective probabilities 客观概率 通常不会因人而异的概率; 包括先验概率和客观概率。

one-sided hypothesis test (or one-tailed hypothesis test) 单边假设检验 (或单侧假设检验) 假设检验的一种, 原假设只有在证据表明总体参数大于 (小于) 数值 θ_0 时才被拒绝。其备择假设同样只有单边。

opportunity cost 机会成本 投资者由于选择特定投资行为而放弃的价值; 这一价值是其他选择中最大的价值。

opportunity set 机会集 可投资资产的集合。

optimizer 最优化程序 用来求解在给定期望收益率水平下具有最小风险的资产组合权重的一段计算机程序或一份电子表格。

ordinal scale 顺序尺度 一种将数据分类的度量尺度, 数据类别按照某种特征进行排序。

ordinary annuity 普通年金 首笔现金流支付发生在当前一个期间之后的年金。

ordinary least squares, OLS 最小二乘法 一种基于最小化回归模型残差平方和准则的估计方法。

orthogonal 正交 不相关; 成直角的。

outcome 结果 随机变量的可能取值。

outliers 异常值 在样本中极端处 (过小或者过大) 的小部分观测值。

out-of-sample forecast errors 样本外预测误差 在用来拟合模型的样本期间以外, 时间序列的实际值和预测值之间的差。

out-of-sample test 样本外检验 用构建策略或模型的时间区间以外的样本对策略或模型进行检验。

paired comparisons test 配对比较检验 基于来自相互依赖的样本的配对的观测值, 为考察总体间差异而进行的统计检验。

paired observations 配对观测值 相互依赖的观测值。

pairs arbitrage trade 配对套利交易 在两个密

- 切相关的股票中的交易，涉及卖空其中一个股票和买入另一个股票。
- panel data **面板数据** 对于多个观测单位的某个特征在不同时间上的观测值。
- parameter **参数** 由总体数据计算得出或者用来描述总体数据特性的描述性度量指标，一般用希腊字母表示。
- parameter instability **参数不稳定性** 总体回归参数随着时间变化的问题或情况。
- parametric test **参数检验** 任何关于参数或者参数的正确性检验，该检验根据关于产生样本的总体的假设进行。
- partial regression coefficients or partial slope coefficients **偏回归系数或者偏斜率系数** 多元回归的斜率系数。
- percentiles **百分位数** 将分布划分为 100 等份的分位数。
- perfect collinearity **完全共线性** 两个或者两个以上自变量或自变量组合间存在精确的线性关系。
- performance appraisal **业绩评价** 对于风险调整后业绩的评价；对于投资技巧的评价。
- performance measurement **业绩度量** 以合理和一致的方式计算收益率。
- periodic rate **期间利率** 每期的报价利率；名义年利率除以一年中复利期间数目。
- permutation **排列** 一个有序的序列。
- perpetuity **永续年金** 永久性的年金，或者一组首笔现金流发生在一期之后的、无穷期的等额现金流序列。
- platykurtic **低峰** 描述比正态分布峰度低的分布。
- point estimate **点估计** 关于一个未知量（如总体参数）的单个数值估计。
- pooled estimate **合并估计** 将来自两个或两个以上样本的观测值联合（合并）起来进行的参数估计。
- population **总体** 指定集合中的所有元素。
- population mean **总体均值** 总体的算术平均值；所有观测值或总体中数据的算术平均值。
- population standard deviation **总体标准差** 度量总体离散程度的指标，具有与观测值相同的度量单位，通过计算总体方差的正平方根得到。
- population variance **总体方差** 度量总体离散程度的指标，通过计算对总体均值的偏差的平方的均值得到。
- portfolio performance attribution **投资组合业绩归因** 根据来自不同风险源的贡献对投资组合业绩进行分析。
- portfolio possibilities curve **投资组合可能曲线** 可以用两种资产构成的所有投资组合的期望收益和风险的图形表示。
- positive serial correlation **正序列相关** 一种序列相关关系，在该关系中一个观测值出现正误差会增加另一个观测值出现正误差的可能性，一个观测值出现负误差会增加另一个观测值出现负误差的可能性。
- posterior probability **后验概率** 反映了新信息的或者是在新信息到达后得到的更新的概率。
- power of a test **检验的势** 正确地拒绝原假设的概率，即当原假设错误时，否定原假设的概率。
- present value **现值** 未来现金流在当前的价值（或折现价值）。
- price relative **价比** 期末价格与期初价格之比；它等于 1 加上该资产的持有期收益率。
- priced risk **定价风险** 投资者需要有额外回报才会承担的风险。
- principal **本金** 最初投资于项目或者金融工具中的资金数量；在到期日所要支付的面值。
- prior probabilities **先验概率** 反映在新信息到达之前对于事件发生可能性的信念的概率。
- probability **概率** 用来描述某事件发生可能性大小的 0 到 1 之间的一个数字。
- probability density function **概率密度函数** 取非负值的函数，概率可以用这个函数所绘制曲线下方的面积来表示。

probability distribution **概率分布** 描述一个随机变量可能出现结果的概率情况的分布。

probability function **概率函数** 描述随机变量取某个特定值的概率的函数。

probit model probit **模型** 基于正态分布的、因变量为定性变量的多元回归模型。

pseudo-random numbers **伪随机数** 由随机数生成器所产生的数。

pure discount instruments **纯贴现工具** 将借款金额与还款金额的差额作为利息支付的工具。

pure factor portfolio **纯因素投资组合** 对于所讨论因子的敏感性为1,而对于所有其他因子的敏感性为0的投资组合。

p-Value **P 值** 最小的显著性水平,在此水平上原假设会被拒绝;也被称为边际显著性水平。

qualitative dependent variables **定性因变量** 用来作为因变量而非自变量的虚拟变量。

quantile (or fractile) **分位数 (或分位点)** 一个数值,小于等于该数值的数据占总数据的某个给定比例。

quartiles **四分位数** 将分布划分为4等份的分位数。

quintiles **五分位数** 将分布划分为5等份的分位数。

random number **随机数** 一个来自均匀分布的观测值。

random number generator **随机数生成器** 产生0~1之间均匀分布的随机数的一种算法。

random variable **随机变量** 未来结果不确定的一个量。

random walk **随机游走** 一种时间序列,其某一期的值为上一期的值加上一个不可预测的随机误差。

range **极差** 一个数据集中最大值与最小值之差。

ratio Scales **比率尺度** 一种度量尺度,它不仅拥有间隔尺度的所有性质,而且还具有一个真正的零点作为原点。

real risk-free interest rate **实际无风险利率**

在没有通货膨胀预期的情况下,一个完全无风险的证券的单期利率。

regime **机制** 关于时间序列的概念,产生时间序列的某个特定模型。

regression coefficients **回归系数** 回归方程式中的截距与斜率系数。

rejection point (or critical value) **拒绝点 (临界值)** 计算出的检验统计量与之进行比较从而判断是否拒绝原假设的一个数值。

relative dispersion **相对离散度** 相对于一个参考值或者基准的分散度的大小。

relative frequency **相对频率** 关于分类数据区间的概念,落在某个区间中的观测值的数量除以样本中观测值的总数量。

residual autocorrelations **残差自相关** 样本残差项之间的自相关性。

risk premium **风险溢价** 投资的期望收益率减去无风险利率。

robust **稳健的** 假设受到违背时,相对不受影响的一种性质。

robust standard errors **稳健标准误** 对回归中误差项存在异方差情况进行校正后所得到的回归估计参数的标准误。

root mean squared error, RMSE **均方误差平方根** 平均平方预测误差的平方根;用来比较预测模型样本外的预测效果。

Roy's safety first criterion **罗伊安全第一准则** 一个判断最优投资组合的准则,认为最优投资组合是使组合收益低于某一阈值的概率最小的那个组合。

rule of 72 **72 法则** 该法则认为使得一项投资的收益翻倍所需的年限大约为72除以名义利率所得的值。

safety-first rules **安全第一法则** 投资组合选择法则,关注在某个时间期限投资组合价值跌至低于某个最低可接受水平的风险。

sample **样本** 总体的一个子集。

sample excess kurtosis **样本超额峰度** 度量一个样本分布尖峰超过标准正态分布尖峰程

- 度的指标。
- sample kurtosis **样本峰度** 度量一个样本分布尖峰程度的指标。
- sample mean **样本均值** 样本观测值之和除以样本大小所得的值。
- sample selection bias **样本选择偏差** 根据特定的属性系统性地把总体中的一些个体排除在外而造成的偏差。例如,当由于数据可得性导致某些观测值无法包含在分析中时,就会造成这一偏差。
- sample skewness **样本偏度** 样本分布的不对称程度的度量指标。
- sample standard deviation **样本标准差** 样本方差的正平方根。
- sample statistic or statistic **样本统计量或统计量** 一个从样本计算得到的或者是用来描述样本的量。
- sample variance **样本方差** 一个度量样本分布分散程度的指标,由对样本均值的离差平方和除以样本大小(n)减1得出。
- sampling **抽样** 获取样本的过程。
- sampling distribution **抽样分布** 由从同一个总体中随机抽取的大小相同的样本计算得到的某个统计量的可能的取值的分布。
- sampling error **抽样误差** 统计量的观测值与要估计的量之间的差。
- sampling plan **抽样计划** 用来选择样本的一套规则。
- scatter plot **散点图** 表示两个数据序列的成对观测值的二维图形。
- security market line, SML **证券市场线** 资本资产定价模型的图形。
- semideviation **半离差** 半方差的正平方根。(有时称为半标准差)
- semilogarithmic **半对数的** 描述所构造的一种尺度,它以纵轴上相等的区间表示相等的变化率,以横轴上相等的区间表示相等的变化量。
- semivariance **半方差** 低于均值的数据的平均离差平方。
- serially correlated **序列相关** 关于回归误差的概念,指各个观测值的误差项之间相关。
- sharpe ratio **夏普比率** 超过无风险利率部分的平均收益率与收益率标准差之间的比率;衡量每一单位收益率标准差所赚取的平均超额收益率的指标。
- shortfall risk **超亏风险** 在一定时间期限,投资组合价值跌至低于某个最小可接受水平的风险。
- simple interest **单利** 每期由初始投资所获得的利息;仅对于本金计算出的利息。
- simple random sample **简单随机样本** 从总体中所选取的一个子集,该子集中的元素都是从总体中以相同的概率选择进来的。
- simple random sampling **简单随机抽样** 一种选择样本的方法,这种方法使得所抽取的样本满足简单随机样本的定义。
- simulation trial **模拟试验** 完成所有模拟步骤的全过程。
- skewed **偏的** 不对称。
- skewness **偏度** 对于偏斜(缺乏对称性)的一种定量度量。与偏斜(skew)同义。
- spearman rank correlation coefficient **斯皮尔曼秩相关系数** 对排序数据相关关系的一种度量。
- spurious correlation **伪相关** 一种误导变量间关系的相关性。
- standard deviation **标准差** 方差的正的平方根;一种与原始数据单位相同的度量分散度的指标。
- standard normal distribution (or unit normal distribution) **标准正态分布(或单位正态分布)** 均值(μ)为0、标准差(σ)为1的正态分布。
- standardized beta **标准化贝塔** 在基本因素模型中,将一个股票的属性值减去所有股票的属性值的平均值再除以所有股票属性值的标准差后得到的值。
- standardizing **标准化** 将数据减去均值后再除以标准差的变换。

stated annual interest rate or quoted interest rate **名义年利率或报价利率** 不考虑年度内复利情况的报价利率。

statistic **统计量** 由样本数据计算出的或者是用来描述样本数据的量。

statistical factor models **统计因子模型** 将统计的方法应用到一组历史收益率数据上来确定能最佳解释历史收益率的协方差或方差的投资组合的多因子模型。

statistical inference **统计推断** 从一小组实际观测到的数据来对一大组数据做出预测、估计或者判断；使用样本统计量来推断未知的总体参数的值。

statistically significant **统计上显著的** 表明原假设能够被拒绝的结果；对于回归的一个估计系数来说，经常被理解为指表明相应总体回归系数不为零的结果。

statistics **统计学** 描述、分析数据，并从数据得出结论的一门学科；也包括数值数据的收集。

stratified random sampling **分层随机抽样** 基于一个或者多个分类标准将总体分成几个子总体（层）的一种抽样过程。从每层中抽取简单随机样本，每层抽取的样本大小与该层相对于总体的大小成比例。然后，再将从各层抽取的样本汇集在一起。

stress testing/scenario analysis **压力测试/情景分析** 一组估计在极端不利事件或者情景下的损失的技术。

subjective probability **主观概率** 基于个人的或者主观的判断而得到的概率。

surprise **意外值** 变量的实际值减去它的预测值（或期望值）。

survivorship bias **生存偏差** 由于在检验设计中忽略了因倒闭、兼并或其他原因在数据库中缺失的公司而造成的一种偏差。

systematic factors **系统因素** 影响大量不同资产的平均收益率的因素。

systematic sampling **系统抽样、等距抽样** 从总体中每隔 k 个个体抽取一个样本单元，直

至达到所需样本大小的一种抽样方式。通过这种抽样方式获得的样本可以近似视为随机的。

target semideviation **目标半离差** 目标半方差的正平方根。

target semivariance **目标半方差** 低于目标值的离差平方的平均值。

t-distribution **t-分布** 由单个参数（自由度）确定的一种对称分布。它被大量用来对方差未知的正态分布的均值进行统计推断。

test statistic **检验统计量** 基于样本数据计算得到的一个量，根据它的值来判断是否拒绝原假设。

time series **时间序列** 在不同时间段对某个变量的一组观测值。

time value of money **货币的时间价值** 关于不同时点现金流之间等价换算关系的原理。

time-period bias **时期偏差** 使用时间序列样本时，统计结果可能会对样本的起始时点和终止时点的选取敏感，从而造成偏差。

time-series data **时间序列数据** 一个变量在不同时点上的观测值。

time-weighted rate of return **时间加权收益率**
在给定量测区间内，一单位货币投资于某个投资组合所获得的复合增长率；是对于投资业绩的一种度量，其对于择时和投资组合中投资额的增减变动不敏感。

total probability rule **全概率公式** 用某事件的以一组互斥且穷举的情景为条件的概率来解释该事件发生的无条件概率的一种规则。

total probability rule for expected value **期望值的全概率公式** 用某随机变量的以一组互斥且穷举的情景为条件的期望值来解释该随机变量的期望值的一种规则。

tracking error **跟踪误差** 同跟踪风险（tracking risk）和积极风险（active risk）。亦指投资组合的总收益（包括各种费用）减去基准的总收益。

tracking portfolio **跟踪投资组合** 一个具有特定因子敏感性的投资组合，通常复制一个基

- 准指数或其他投资组合。
- tracking risk **跟踪风险** 资产组合与其基准资产组合的收益之差的标准差。同积极风险 (active risk)。
- tree diagram **树状图** 由各个节点产生分支来表示互斥的随机事件或互斥的决策的图形。
- trend **趋势** (时间序列) 具有特定方向性的长期变动模式。
- trimmed mean **截尾均值** 按照一个给定的很小的比例去掉若干最小和最大的观测值之后计算得到的均值。
- t-Test **t 检验** 使用服从 t 分布的统计量 (t 统计量) 进行的假设检验。
- two-sided hypothesis test (or two-tailed hypothesis test) **双边假设检验** (或者是双尾假设检验) 一种假设检验, 当证据表明总体参数小于或者大于某个假定的值时就拒绝原假设而接受备择假设。
- type I error **第一类错误** 拒绝正确的原假设而造成的错误。
- type II error **第二类错误** 没有拒绝错误的原假设而造成的错误。
- unbiasedness **无偏性** 没有偏差。这是想要的估计量的性质。一个无偏估计量是指其期望值 (抽样分布的均值) 等于所要估计的参数的值。
- unconditional heteroskedasticity **无条件异方差** 与回归中自变量的值不相关的误差项的异方差性。
- unconditional probability (or marginal probability) **无条件概率 (或边缘概率)** 一个事件的不以其他事件为条件的概率。
- unit root **单位根** 不具有协方差平稳性质的时间序列被称为具有单位根。
- univariate distribution **单变量分布** 单个随机变量的概率分布。
- up transition probability **向上转移概率** 资产价值上升的概率。
- value at Risk, VAR **在险价值** 在给定的概率水平下, 某个指定期间的以货币计量的最小的期望损失值。
- variance **方差** 与随机变量期望值的离差平方的期望值 (概率加权平均值)
- volatility **波动率** 如在期权定价中所使用的, 标的资产的连续复利收益率的标准差。
- weighted mean **加权平均** 对于每个观测值用其相对重要性指标予以加权所算得的平均值。
- weighted-average cost of capital, WACC **加权平均资本成本** 公司普通股、优先股和长期债券的税后要求收益率的加权平均值, 其中的权重为公司目标资本结构中各个融资来源所占的比重。
- white-corrected standard errors **怀特修正的标准误** 同稳健的标准误 (robust standard errors)
- winsorized mean **缩尾均值** 将给定比例的最小值均取值为某个指定的小的数值, 将给定比例的最大值均取值为某个指定的大的数值, 之后计算得到的均值。
- working capital management **营运资本管理** 对于企业短期资产 (如存货) 和短期负债 (如欠供应商的款项) 的管理。



参考文献

- Ackerman, Carl, Richard McEnally, and David Ravenscraft. 1999. "The Performance of Hedge Funds: Risk, Return, and Incentives." *Journal of Finance*. Vol. 54, No. 3: 833–874.
- Altman, Edward I. 1968. "Financial Ratios, Discriminant Analysis and the Prediction of Corporate Bankruptcy." *Journal of Finance*. Vol. 23: 589–699.
- Altman, Edward I., and Vellore M. Kishore. 1996. "Almost Everything You Wanted to Know about Recoveries on Defaulted Bonds." *Financial Analysts Journal*. Vol. 52, No. 6: 57–63.
- Altman, Edward I., R. Haldeman, and P. Narayanan. 1977. "Zeta Analysis: A New Model to Identify Bankruptcy Risk of Corporations." *Journal of Banking and Finance*. Vol. 1: 29–54.
- Ambachtsheer, Keith, Ronald Capelle, and Tom Scheibelhut. 1998. "Improving Pension Fund Performance." *Financial Analysts Journal*. Vol. 54, No. 6: 15–21.
- American Association of Individual Investors. 2003. *The Individual Investor's Guide to the Top Mutual Funds*, 22nd edition. Chicago: AAI.
- Amihud, Yakov, and Kefei Li. 2002. "The Declining Information Content of Dividend Announcements and the Effect of Institutional Holdings." New York University Stern School of Business Working Paper No. SC-AM-02-11.
- Amin, Gaurav S., and Harry M. Kat. 2003. "Stocks, Bonds, and Hedge Funds." *Journal of Portfolio Management*. Vol. 29, No. 4: 113–120.
- Baks, Klaas P., Andres Metrick, and Jessica Wachter. 2001. "Should Investors Avoid All Actively Managed Mutual Funds?" *Journal of Finance*. Vol. 56: 45–85.
- Barr, G.D.I., and Kantor, B.S. 1999. "Price-Earnings Ratios on the Johannesburg Stock Exchange—Are They a Guide to Value?" *SA Journal of Accounting Research*. Vol. 13, No. 1: 1–23.
- Bauman, W. Scott, C. Mitchell Conover, and Robert E. Miller. 1998. "Growth versus Value and Large-Cap versus Small-Cap Stocks in International Markets." *Financial Analysts Journal*. Vol. 54, No. 2: 75–89.
- Black, Fischer. 1972. "Capital Market Equilibrium with Restricted Borrowing." *Journal of Business*. Vol. 45, No. 3: 444–455.
- . 1993. "Estimating Expected Return." *Financial Analysts Journal*. Vol. 49, No. 5: 36–38.
- Black, Fischer, and Robert Litterman. 1991. "Asset Allocation: Combining Investor Views with Market Equilibrium." *Journal of Fixed Income*. Vol. 1, No. 2: 7–18.
- . 1992. "Global Portfolio Optimization." *Financial Analysts Journal*. Vol. 48, No. 5: 28–43.
- Block, Stanley B. 1999. "A Study of Financial Analysts: Practice and Theory." *Financial Analysts Journal*. Vol. 55, No. 4: 86–95.
- Blume, Marshall. 1984. "The Use of 'Alphas' to Improve Performance." *Journal of Portfolio Management*. Vol. 11: 86–92.
- Bodie, Zvi, Alex Kane, and Alan J. Marcus. 2001. *Investments*, 5th edition. Boston: Irwin/McGraw-Hill.
- Bowerman, Bruce L., and Richard T. O'Connell. 1997. *Applied Statistics: Improving Business Practices*. Chicago: Irwin.
- Brealey, Richard A., and Stewart C. Myers. 2003. *Principles of Corporate Finance*. 7th ed. New York: McGraw-Hill.
- Breusch, T., and A. Pagan. 1979. "A Simple Test for Heteroscedasticity and Random Coefficient Variation." *Econometrica*. Vol. 47: 1287–1294.

- Brown, Stephen, William Goetzmann, and Stephen Ross. 1995. "Survival." *Journal of Finance*. Vol. 50: 853–873.
- Buetow, Gerald W., Jr., Robert R. Johnson, and David E. Runkle. 2000. "The Inconsistency of Return-Based Style Analysis." *Journal of Portfolio Management*. Vol. 26, No. 3: 61–77.
- Buffett, Warren. 1993. *Berkshire Hathaway Chairman's Letter*. Retrieved May 25, 2001, from the World Wide Web: www.berkshirehathaway.com/letters/1983.html.
- Burmeister, Edwin, and Marjorie B. McElroy. 1988. "Joint Estimation of Factor Sensitivities and Risk Premia for the Arbitrage Pricing Theory." *Journal of Finance*. Vol. 43, No. 3: 721–733.
- Burmeister, Edwin, Richard Roll, and Stephen A. Ross. 1994. "A Practitioner's Guide to Arbitrage Pricing Theory." In *A Practitioner's Guide to Factor Models*. Charlottesville, VA: Research Foundation of the Institute of Chartered Financial Analysts.
- Campbell, John Y., Andrew W. Lo, and A. Craig MacKinlay. 1997. *The Econometrics of Financial Markets*. Princeton, NJ: Princeton University Press.
- Campbell, John Y., Martin Lettau, Burton G. Malkiel, and Yexiao Xu. 2001. "Have Individual Stocks Become More Volatile? An Empirical Exploration of Idiosyncratic Risk." *Journal of Finance*. Vol. 56, No. 1: 1–43.
- Campbell, Stephen K. 1974. *Flaws and Fallacies in Statistical Thinking*. Englewood Cliffs, NJ: Prentice-Hall.
- Carhart, Mark M. 1997. "On Persistence in Mutual Fund Performance." *Journal of Finance*. Vol. 52, No. 1: 57–82.
- Chan, Louis K.C., Jason Karceski, and Josef Lakonishok. 1999. "On Portfolio Optimization: Forecasting Covariances and Choosing the Risk Model." *Review of Financial Studies*. Vol. 12, No. 5: 937–974.
- Chance, Don M. 2003. *Analysis of Derivatives for the CFA® Program*. Charlottesville, VA: AIMR.
- Chen, Nai-fu, Richard Roll, and Stephen Ross. 1986. "Economic Forces and the Stock Market." *Journal of Business*. Vol. 59, No. 3: 383–403.
- Chow, George. 1995. "Portfolio Selection Based on Return, Risk, and Relative Performance." *Financial Analysts Journal*. Vol. 51, No. 2: 54–60.
- Chow, George, Eric Jacquier, Mark Kritzman, and Kenneth Lowry. 1999. "Optimal Portfolios in Good Times and Bad." *Financial Analysts Journal*. Vol. 55, No. 3: 65–73.
- Chua, Jess H., Richard S. Woodward, and Eric C. To. 1987. "Potential Gains from Stock Market Timing in Canada." *Financial Analysts Journal*. Vol. 43, No. 5: 50–56.
- Clarke, Roger G., and Mark P. Kritzman. 1996. *Currency Management: Concepts and Practices*. Charlottesville, VA: Research Foundation of AIMR.
- Cochrane, John H. 1999a. "New Facts in Finance." *Economic Perspectives*. Federal Reserve Bank of Chicago. Vol. 23, No. 3: 36–58. (Revision of NBER Working Paper 7169.)
- . 1999b. "Portfolio Advice for a Multifactor World." *Economic Perspectives*. Federal Reserve Bank of Chicago. Vol. 23, No. 3: 59–78. (Revision of NBER Working Paper 7170.)
- Connor, Gregory. 1995. "The Three Types of Factor Models: A Comparison of Their Explanatory Power." *Financial Analysts Journal*. Vol. 51, No. 3: 42–46.
- Cox, Jonathan, Stephen Ross, and Mark Rubinstein. 1979. "Options Pricing: A Simplified Approach." *Journal of Financial Economics*. Vol. 7: 229–263.
- Daniel, Wayne W., and James C. Terrell. 1986. *Business Statistics, Basic Concepts and Methodology*, 4th edition. Boston: Houghton Mifflin.
- . 1995. *Business Statistics for Management & Economics*, 7th edition. Boston: Houghton-Mifflin.
- Davidson, Russell, and James G. MacKinnon. 1993. *Estimation and Inference in Econometrics*. New York: Oxford University Press.
- Dickey, David A., and Wayne A. Fuller. 1979. "Distribution of the Estimators for Autoregressive Time Series with a Unit Root." *Journal of the American Statistical Association*. Vol. 74: 427–431.
- Diebold, Francis X. 2004. *Elements of Forecasting*, 3rd edition. Cincinnati, OH: South-Western.
- Dimson, Elroy, Paul Marsh, and Mike Staunton. 2002. *Triumph of the Optimists: 101 Years of Global Investment Returns*. Princeton, NJ: Princeton University Press.

- Durbin, J., and G.S. Watson. 1951. "Testing for Serial Correlation in Least Squares Regression, II." *Biometrika*. Vol. 38: 159–178.
- Dybvig, Philip H., and Stephen A. Ross. 1985a. "Differential Information and Performance Measurement Using a Security Market Line." *Journal of Finance*. Vol. 40, No. 2: 383–399.
- . 1985b. "The Analytics of Performance Measurement Using a Security Market Line." *Journal of Finance*. Vol. 40, No. 2: 401–416.
- Edwards, Franklin R. 1988. "Does Futures Trading Increase Stock Market Volatility?" *Financial Analysts Journal*. Vol. 44, No. 1: 63–69.
- Elton, Edwin J., Martin J. Gruber, and Joel Rentzler. 1987. "Professionally Managed Publicly Traded Commodity Funds." *Journal of Business*. Vol. 60, No. 2: 175–199.
- . 1990. "The Performance of Publicly Offered Commodity Funds." *Financial Analysts Journal*. Vol. 46, No. 4: 23–30.
- Elton, Edwin J., Martin J. Gruber, Stephen J. Brown, and William N. Goetzmann. 2003. *Modern Portfolio Theory and Investment Analysis*, 6th edition. New York: Wiley.
- Engle, Robert F. 1982. "Autoregressive Conditional Heteroscedasticity with Estimates of the Variance of United Kingdom Inflation." *Econometrica*. Vol. 50, No. 4: 987–1007.
- Engle, Robert F., and Clive W.J. Granger. 1987. "Co-Integration and Error Correction: Representation, Estimation, and Testing." *Econometrica*. Vol. 55, No. 2: 251–276.
- Estrada, Javier. 2003. "Mean-Semivariance Behavior: A Note." *Finance Letters*. Vol. 1, Issue 1: 9–14.
- Fabozzi, Frank J. 2004a. *Fixed Income Analysis for the Chartered Financial Analyst® Program*, 2nd edition. New Hope, PA: Frank J. Fabozzi Associates.
- . 2004b. *Fixed Income Readings for the Chartered Financial Analyst® Program*, 2nd edition. New Hope, PA: Frank J. Fabozzi Associates.
- Fama, Eugene. 1976. *Foundations of Finance*. New York: Basic Books.
- Fama, Eugene F., and Kenneth R. French. 1992. "The Cross-Section of Expected Stock Returns." *Journal of Finance*. Vol. 47, No. 2: 427–466.
- . 1993. "Common Risk Factors in the Returns on Stocks and Bonds." *Journal of Financial Economics*. Vol. 33, No. 1: 3–56.
- . 1996. "Multifactor Explanations of Asset Pricing Anomalies." *Journal of Finance*, Vol. 51, No. 1: 55–84.
- Feller, William. 1957. *An Introduction to Probability Theory and Its Applications, Vol. I*, 2nd edition. New York: Wiley.
- Ferguson, Robert. 1993. "Some Formulas for Evaluating Two Popular Option Strategies." *Financial Analysts Journal*. Vol. 49, No. 5: 71–76.
- Fisher, Lawrence, and James H. Lorie. 1970. "Some Studies of Variability of Returns on Investments in Common Stocks." *Journal of Business*. Vol. 43, No. 2: 99–134.
- Forbes*. 16 September 2003. "The Honor Roll." New York: Forbes Management Co., Inc.
- Freeley, Austin J., and David L. Steinberg. 1999. *Argumentation and Debate: Critical Thinking for Reasoned Decision Making*. Wadsworth Publishing.
- French, Kenneth R., and James M. Poterba. 1991. "Investor Diversification and International Equity Markets." *American Economic Review*. Vol. 81, No. 2: 222–226.
- Freund, John E., and Frank J. Williams. 1977. *Elementary Business Statistics*, 3rd edition. Englewood Cliffs, NJ: Prentice-Hall.
- Fridson, Martin S., and M. Christopher Garman. 1998. "Determinants of Spreads on New High-Yield Bonds." *Financial Analysts Journal*. Vol. 54, No. 2: 28–39.
- Fung, William, and David A. Hsieh. 2002. "Hedge-Fund Benchmarks: Information Content and Biases." *Financial Analysts Journal*. Vol. 58, No. 1: 22–34.
- Garbade, Kenneth. 1982. *Securities Markets*. New York: McGraw-Hill.
- Goetzmann, William N., and Philippe Jorion. 1999. "Re-Emerging Markets." *Journal of Financial and Quantitative Analysis*. Vol. 34, No. 1: 1–32.
- Goldberger, Arthur S. 1998. *Introductory Econometrics*. Cambridge, MA: Harvard University Press.

- Granger, Clive W.J., and Paul Newbold. 1974. "Spurious Regressions in Econometrics." *Journal of Econometrics*. Vol. 2: 111–120.
- Greene, William H. 2003. *Econometric Analysis*, 5th edition. Upper Saddle River, NJ: Prentice-Hall.
- Grinblatt, Mark, and Sheridan Titman. 1998. *Financial Markets and Corporate Strategy*. New York: Irwin/McGraw-Hill.
- Grinold, Richard, and Ronald N. Kahn. 1994. "Multi-Factor Models for Portfolio Risk." In *A Practitioner's Guide to Factor Models*. Charlottesville, VA: Research Foundation of the Institute of Chartered Financial Analysts.
- Gujarati, Damodar N. 2003. *Basic Econometrics*, 4th edition. New York: McGraw-Hill.
- Hamilton, James D. 1994. *Time Series Analysis*. Princeton, NJ: Princeton University Press.
- Hansen, Lars Peter. 1982. "Large Sample Properties of Generalized Method of Moments Estimators." *Econometrica*. Vol. 50, No. 4: 1029–1054.
- Heaney, Richard, Chitoshi Koga, Barry Oliver, and Alfred Tran. 1999. "The Size Effect and Derivative Usage in Japan." Working paper: The Australian National University.
- Henriksson, Roy D., and Robert C. Merton. 1981. "On Market Timing and Investment Performance, II. Statistical Procedures for Evaluating Forecasting Skills." *Journal of Business*. Vol. 54, No. 4: 513–533.
- Hettmansperger, T.P., and J.W. McKean. 1998. *Robust Nonparametric Statistical Methods*. London: Oxford University Press.
- Hillier, Frederick S., and Gerald J. Lieberman. 2000. *Introduction to Operations Research*, 7th edition. New York: McGraw-Hill.
- Hull, John. 2003. *Options, Futures, and Other Derivatives*, 5th edition. Upper Saddle River, NJ: Prentice-Hall.
- Kahn, Ronald N., and Andrew Rudd. 1995. "Does Historical Performance Predict Future Performance?" *Financial Analysts Journal*. Vol. 51, No. 6: 43–52.
- Keane, Michael P., and David E. Runkle. 1990. "Testing the Rationality of Price Forecasts: New Evidence from Panel Data." *American Economic Review*. Vol. 80, No. 4: 714–735.
- . 1998. "Are Financial Analysts' Forecasts of Corporate Profits Rational?" *Journal of Political Economy*. Vol. 106, No. 4: 768–805.
- Kemeny, John G., Arthur Schleifer, Jr., J. Laurie Snell, and Gerald L. Thompson. 1972. *Finite Mathematics with Business Applications*, 2nd edition. Englewood Cliffs, NJ: Prentice-Hall.
- Klemkosky, Robert C., and John D. Martin. 1975. "The Adjustment of Beta Forecasts." *Journal of Finance*. Vol. 30, No. 4: 1123–1128.
- Kmenta, Jan. 1986. *Elements of Econometrics*, 2nd edition. New York: Macmillan.
- Kolb, Robert W., Gerald D. Gay, and William C. Hunter. 1985. "Liquidity Requirements for Financial Futures Investments." *Financial Analysts Journal*. Vol. 41, No. 3: 60–68.
- Kon, Stanley J. 1984. "Models of Stock Returns—A Comparison." *Journal of Finance*. Vol. 39: 147–165.
- Kool, Clemens J.M. 2000. "International Bond Markets and the Introduction of the Euro." *Review of the Federal Reserve Bank of St. Louis*. Vol. 82, No. 5: 41–56.
- Kothari, S.P., Jay Shanken, and Richard G. Sloan. 1995. "Another Look at the Cross-Section of Expected Stock Returns." *Journal of Finance*. Vol. 50, No. 1: 185–224.
- Ledoit, Olivier, and Michael Wolf. 2004. "Honey, I Shrunk the Sample Covariance Matrix." *Journal of Portfolio Management*. Vol. 31, No. 1.
- Leibowitz, Martin L., and Roy D. Henriksson. 1989. "Portfolio Optimization with Shortfall Constraints: A Confidence-Limit Approach to Managing Downside Risk." *Financial Analysts Journal*. Vol. 45, No. 2: 34–41.
- Leinweber, David. 1997. *Stupid Data Mining Tricks: Over-Fitting the S&P 500*. Monograph. Pasadena, CA: First Quadrant.
- Liang, Bing. 1999. "On the Performance of Hedge Funds." *Financial Analysts Journal*. Vol. 55, No. 4: 72–85.

- Linsmeier, Thomas J., and Neil D. Pearson. 2000. "Value at Risk." *Financial Analysts Journal*. Vol. 56, No. 2: 47–67.
- Lo, Andrew W. 1999. "The Three P's of Total Risk Management." *Financial Analysts Journal*. Vol. 55, No. 1: 13–26.
- Lo, Andrew W., and A. Craig MacKinlay. 1990. "Data Snooping Biases in Tests of Financial Asset Pricing Models." *Review of Financial Studies*. Vol. 3: 175–208.
- Luenberger, David G. 1998. *Investment Science*. New York: Oxford University Press.
- MacKinlay, A. Craig, and Matthew P. Richardson. 1991. "Using Generalized Methods of Moments to Test Mean–Variance Efficiency." *Journal of Finance*. Vol. 46, No. 2: 511–527.
- Mankiw, N. Gregory. 2000. *Macroeconomics*, 4th edition. New York: Worth Publishers.
- Mauboussin, Michael J., and Shelby Heard. 1996. *Packaged Food Industry Rate of Return Analysis*. New York: Credit Suisse First Boston.
- . 1975. "Selecting Economic Hypotheses by Goodness of Fit." *Economic Journal*. Vol. 85, Issue 340: 877–883.
- Mayer, Thomas. 1980. "Economics as a Hard Science: Realistic Goal or Wishful Thinking?" *Economic Inquiry*. Vol. 18, No. 2: 165–178.
- McQueen, Grant, and Steven Thorley. 1999. "Mining Fools Gold." *Financial Analysts Journal*. Vol. 55, No. 2: 61–72.
- McQueen, Grant, Kay Shields, and Steven R. Thorley. 1997. "Does the 'Dow-10 Investment Strategy' Beat the Dow Statistically and Economically?" *Financial Analysts Journal*. Vol. 53, No. 4: 66–72.
- Michaud, Richard O. 1998. *Efficient Asset Management*. Boston: Harvard Business School Press.
- Moore, David S., and George P. McCabe. 1998. *Introduction to the Practice of Statistics*, 3rd edition. New York: W.H. Freeman.
- Nayar, Nandkumar, and Michael S. Rozeff. 1994. "Ratings, Commercial Paper, and Equity Returns." *Journal of Finance*. Vol. 49, No. 4: 1431–1449.
- Nelson, David C., Robert B. Moskow, Tiffany Lee, and Gregg Valentine. 2003. *Food Investor's Handbook*. New York: Credit Suisse First Boston.
- Oppenheimer, Henry R., and Sanjiv Sabherwal. 2003. "The Competitive Effects of U.S. Decimalization: Evidence from the U.S.-Listed Canadian Stocks." *Journal of Banking and Finance*. Vol. 27, No. 9.
- Pindyck, Robert S., and Daniel L. Rubinfeld. 1998. *Econometric Models and Economic Forecasts*, 4th edition. Boston: Irwin/McGraw-Hill.
- Ramsey, Frank P. 1931. "Truth and Probability." In *The Foundations of Mathematics and Other Logical Essays*, edited by R.B. Braithwaite. London: Routledge and Keegan Paul.
- Reilly, Frank K., and Keith C. Brown. 2003. *Investment Analysis and Portfolio Management*, 7th edition. Mason, OH: South-Western.
- Roll, Richard, and Stephen A. Ross. 2001. *What Is the Arbitrage Pricing Theory?* Retrieved May 25, 2001, from the World Wide Web: www.rollross.com/apt.html.
- Ross, Sheldon M. 1997. *A First Course in Probability*, 5th edition. Englewood, NJ: Prentice-Hall.
- Ross, Stephen A. 1976. "The Arbitrage Theory of Capital Asset Pricing." *Journal of Economic Theory*. December: 341–360.
- Roy, A.D. 1952. "Safety-First and the Holding of Assets." *Econometrica*. Vol. 20: 431–439.
- Sharpe, William F. 1988. "Determining a Fund's Effective Asset Mix." *Investment Management Review*. November/December: 59–69.
- . 1992. "Asset Allocation: Management Style and Performance Measurement." *Journal of Portfolio Management*. Vol. 18, No. 2: 7–19.
- . 1994. "The Sharpe Ratio." *Journal of Portfolio Management*. Vol. 21, No. 1: 49–59.
- Sharpe, William F., Gordon J. Alexander, and Jeffery V. Bailey. 1999. *Investments*, 6th edition. Upper Saddle River, NJ: Prentice-Hall.
- Shumway, Tyler, and Vincent A. Warther. 1999. "The Delisting Bias in CRSP's Nasdaq Data and Its Implications for the Size Effect." *Journal of Finance*. Vol. 54, No. 6: 2361–2379.
- Siegel, Jeremy J. 1998. *Stocks for the Long Run*, 2nd edition. New York: McGraw-Hill.
- Siegel, Sidney. 1956. *Nonparametric Statistics for the Behavioral Sciences*. New York: McGraw-Hill.

- Sorenson, Eric H., Vele Samak, and Keith Miller. 1998. "U.S. Equity Risk Attribute Model (RAM): A New Look at a Classic Model." *Quantitative Research*. Salomon Smith Barney.
- Standard Life Investment, Inc. 2003. "Pooled Investment Funds Performance Review: First Quarter 2003."
- Stewart, G. Bennett. 1991. *The Quest for Value*. New York: HarperBusiness.
- Stoll, Hans R. 1978. "The Pricing of Security Dealer Services: An Empirical Study of Nasdaq Stocks." *Journal of Finance*. Vol. 33, No. 4: 1153–1172.
- Stowe, John D., Thomas R. Robinson, Jerald E. Pinto, and Dennis W. McLeavey. 2002. *Analysis of Equity Investments: Valuation*. Charlottesville, VA: AIMR.
- Stuart, Alan, and J. Keith Ord. 1994. *Kendall's Advanced Theory of Statistics, Vol. 1: Distribution Theory*, 6th edition. London: Arnold.
- Swensen, David F. 2000. *Pioneering Portfolio Management: An Unconventional Approach to Institutional Investment*. New York: Free Press.
- Treynor, Jack L., and Kay Mazuy. 1966. "Can Mutual Funds Outguess the Market?" *Harvard Business Review*. Vol. 44: 131–136.
- White, Gerald I., Ashwinpaul C. Sondhi, and Dov Fried. 2003. *The Analysis and Use of Financial Statements*, 3rd edition. New York: Wiley.
- White, Halbert. 1980. "A Heteroskedasticity-Consistent Covariance Matrix Estimator and a Direct Test for Heteroskedasticity." *Econometrica*. Vol. 48, No. 4: 817–838.



CFA项目介绍

特许金融分析师（CFA）证书是衡量投资专业人员能力和诚信的全球认证标准。为了获得CFA证书，考生必须顺利通过CFA项目，它是等同于硕士学位程度的自学项目，包含了投资专业人员所必备的专业技能的广泛课程。

以实践课程为基础，CFA项目聚焦由专业人员确认的投资决策过程必备知识，并通过对全球CFA执业分析师进行定期、集中的调查来更新该知识体系。全部课程总共围绕十个主题，内容从股权和固定收益分析到投资组合管理再到公司金融，都强调了专业实践中的道德问题。CFA项目课程以严谨而全面出名，它介绍了适用于所有市场的原理，因此获得CFA证书的专业人员都会具有完备的全球性投资视角，并对全球市场有深刻的理解。

www.cfainstitute.org

作者简介



理查德 A. 德弗斯科，特许金融分析师，内布拉斯加大学林肯分校（Nebraska - Lincoln, UNL）金融学副教授。他于 1999 年获得特许金融分析师的执照。德弗斯科是奥马哈·林肯金融分析师协会成员，同时是 CFA 协会的委员会成员。德弗斯科博士主要教授投资学，并协调内布拉斯加州基金（一家由内布拉斯加大学林肯分校学生管理的投资基金）的日常运营。他的论文（主要是关于金融学领域的论文）曾在许多期刊上发表。德弗斯科博士在罗德岛大学获得了管理科学的学士学位，并在田纳西·诺克斯维尔大学获得了金融学博士学位。

丹尼斯 W·麦克利维，特许金融分析师，是 CFA 协会职业发展部的主管。他于 1990 年获得了特许金融分析师的执照。在 2000 年加入 CFA 协会之前，他曾在多个不同机构的委员会中任职。最近，他与布鲁诺·索尼克合作出版了《国际投资学》（第 5 版）。麦克利维同时是 CFA 协会出版的教科书《定量投资分析》和《投资组合管理：动态过程》的作者之一，另外，他还是有关生产计划和存货控制以及管理决策中的运筹学的作者之一。在麦克利维 25 年的学术生涯中，他曾任教于西安大略大学、康乃迪克大学、罗德岛大学（在这里，他创建了一家由学生管理的基金）以及巴布森学院。他还是纽约股票交易所的一名仲裁员。在西安大略大学获得经济学学士学位之后，麦克利维又在印第安纳大学获得了生产管理和工业工程博士学位。

杰拉尔德 E. 平托，特许金融分析师，是 CFA 协会 CFA 和 CIPM 项目部的主任。他在此之前的工作是纽约市 TRM 服务部的负责人。在从事这一职务之时，他曾为有限责任公司、基金会以及合伙企业提供投资计划、投资组合分析以及定量投资分析方面的咨询服务。平托博士在纽约市投资和银行业工作卸任之后，曾在纽约大学斯特恩商学院教授过金融学。平托是教科书《股权资产估值》的作者之一。他在布鲁克学院获得了 MBA 学位，并在斯特恩商学院获得了金融学博士学位。

戴维 E. 朗克尔，特许金融分析师，是美国银行贾弗瑞公司的副主席和研究部经理。他自 1989 年起就是明尼苏达大学卡尔森管理学院的兼职教授，在那里教授权益证券分析课程。在加入美国银行贾弗瑞公司之前，朗克尔博士是明尼阿波利斯联邦储备银行的一位研究员。他已经发表了超过 20 篇学术论文，并获得了包括表彰杰出教师的温斯顿奖（布朗大学）、表彰在注册公共会计师考试中杰出表现的伊利亚 - 瓦特 - 塞尔斯奖以及在《商业周刊》对于最好商学院杰出教授的四星评级在内的许多奖项。朗克尔博士是明尼苏达注册公共会计师协会的成员，还是 CFA 金融会计政策委员会的成员。他在卡尔顿学院以优异的成绩获得了经济学学士学位，并在麻省理工学院获得了经济学博士学位。



译者后记

特许金融分析师 (Chartered Financial Analyst, CFA) 是由美国 CFA 协会颁发的、代表全球投资业的高标准与权威性的资格称号。自 1962 年设立 CFA 课程以来, CFA 协会为全球投资业在道德操守、专业标准及知识体系等方面设立了全球性的规范与标准。随着 CFA 课程的不断完善以及在全球范围内的大力推广, 这些规范与标准获得了人们广泛的认知与认可。

近年来, 随着我国金融市场的不断开放与发展, 对于金融人才特别是具有全球视野的高端金融人才的需求与日俱增。CFA 特许资格认证持证人在金融业职场上享有明显的竞争优势, CFA 考试课程正在为我国培养具有与国际接轨的专业知识和道德操守的金融从业人员发挥着积极的作用。本书是为 CFA 资格考试中非常重要的一项内容——定量投资分析而编写的一本教材。该教材旨在培养学员的定量分析能力, 使他们能够从科学、客观的角度对金融市场进行认知与评价, 从而做出有效的投资决策。不同于传统的大学金融专业的教材, 本书在介绍每个概念的时候, 都会辅之以非常贴近现实的例子, 对原本较为抽象的概念进行具体的说明, 这使读者能更加充分地理解概念, 并且了解每个概念在现实中是如何得到应用的。这种写作方法对于以实战为目标的投资行业从业人员来说, 可能更加具有实际意义。通过这样的方式所介绍的内容, 也会更容易为读者所接受。本书的作者均为资深的业界与学术界专业人士, 在定量投资分析领域具有非常扎实的理论功底与丰富的实践经验, 这使得本书的内容选择、结构安排和逻辑体系能够很好地反映当代投资业界所采用的定量投资分析方法的最新相关内容。在书中, 作者以系统而简洁明快的语言将业界在解决实际问题中所使用的定量方法介绍给读者, 相信每一位读者都能够在阅读本书的过程中获益匪浅。

阅读本书并不需要太多的基础知识, 作者从基本的概念开始介绍, 循序渐进, 在介绍理论方法的时候伴随着生动具体的例子, 使读者逐步掌握现代金融业先进的定量投资分析技术。如果您是一位金融界的新手, 那么您可以从本书中了解定量投资分析的基本原理 (如货币的时间价值、资产价值的评估方法等), 学会使用统计学知识对市场进行分析的方法 (如市场收益率特征的描述方法、对于假设的检验方法、变量间相关性分析、回归分析、时间序列分析方法等), 以及熟悉行业中投资组合管理的基本分析模型与方法。如果您是一位在金融界工作多年的专业人士, 您同样可以从本书中了解到定量投资分析领域与实务相关的最新发展成果。所以, 本书对于所有对金融投资有兴趣的从业者来说都是一本有参考价值的好书。

我们有幸受机械工业出版社华章公司之托将此书译成中文, 希望中文版的出版能够帮助我国更多读者了解定量投资分析的基础知识, 同时也希望我国的金融界能够通过学习国际先进的定量投资分析的理论和方法, 将相关知识应用到我国本土金融市场的定量分析之中, 以提高我国金融业整体的投资分析和投资组合管理水平。

本书的翻译工作伴随我们度过了 2012 年的春节, 其中的繁忙与辛苦融入书中的字里行间, 也将成为我们的一段难忘的记忆。限于时间和自身水平, 本书的翻译难免会有所疏漏甚至谬误, 恳请广大读者不吝指正, 我们不胜感激。

劳兰珺 王祺

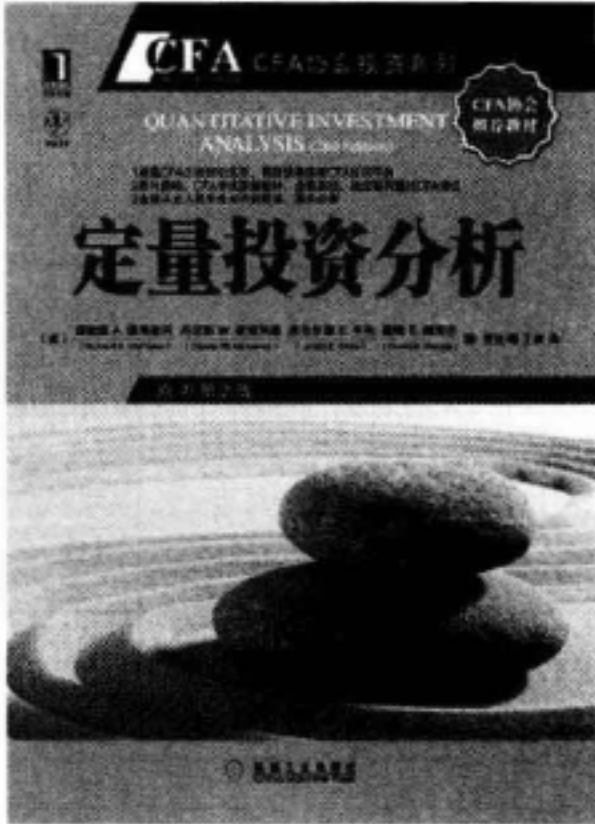
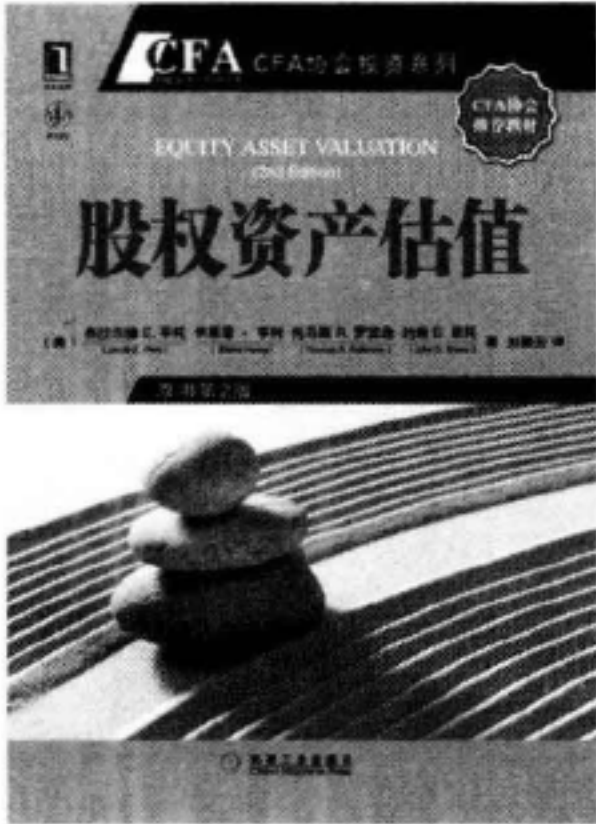
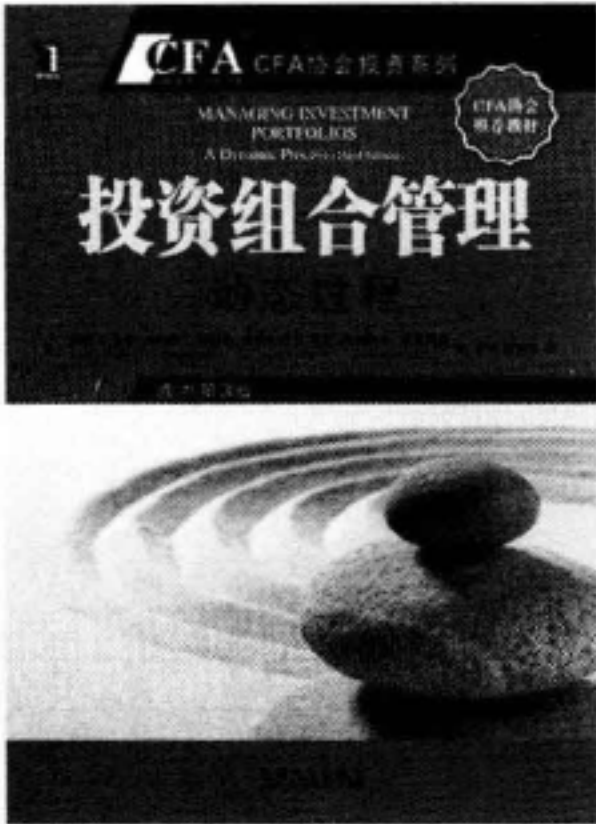
2012 年 6 月于复旦大学

金融学

课程名称	书号	书名、作者及出版时间	版别	定价
行为金融学	978-7-111-34808-5	行为金融（福布斯）（2011年）	外版	62
行为金融学	即将出版	行为金融：心理、决策和市场（阿克特）（2012年）	外版	52
投资银行学	即将出版	投资银行、对冲基金和私募股权投资导论（斯托维尔）（2012年）	外版	68
投资管理	978-7-111-30107-3	投资分析与组合管理（英文版·第9版）（赖利）（2010年）	外版	98
投资管理	978-7-111-22595-9	现代投资组合理论与投资分析（第7版）（埃尔顿）（2007年）	外版	68
投资管理	978-7-111-21204-5	现代投资组合理论与投资分析（英文版·第7版）（埃尔顿）（2007年）	外版	80
商业银行经营管理学	978-7-111-34584-8	商业银行管理（第8版）（罗斯）（2011年）	外版	78
商业银行经营管理学	978-7-111-35539-7	商业银行管理（英文版·第8版）（罗斯）（2011年）	外版	88
金融学（货币银行学）指导或案例	即将出版	货币金融学（第2版）学习指导（米什金）（2011年）	外版	28
金融学（货币银行学）	978-7-111-23245-2	货币、金融体系与经济（英文版·第5版）（哈伯德）（2008年）	外版	65
金融学（货币银行学）	978-7-111-34261-8	货币金融学（第2版）（米什金）（2011年）	外版	75
金融学（货币银行学）	978-7-111-36279-1	货币金融学（克鲁肖）（2011年）	外版	49
金融学（货币银行学）	978-7-111-22274-3	货币银行学：货币、银行业和金融市场（托马斯）（2008年）	外版	78
金融学（货币银行学）	978-7-111-26584-9	金融学导论：市场、投资与财务管理（第13版）（梅利歇尔）（2009年）	外版	75
金融市场学	978-7-111-26909-0	货币、金融市场与金融机构（第10版）（基德维尔）（2009年）	外版	78
金融市场学	978-7-111-26674-7	金融市场学（第10版）（罗斯）（2009年）	外版	79
金融市场学	978-7-111-26790-4	金融市场学（英文版·第10版）（罗斯）（2009年）	外版	88
金融市场学	978-7-111-29552-5	金融市场与机构（第8版）（马杜拉）（2010年）	外版	79
金融市场学	978-7-111-25441-6	金融市场与金融机构（第5版）（米什金）（2008年）	外版	69
金融市场学	即将出版	金融市场与金融机构（第7版）（米什金）（2012年）	外版	69
金融市场学	978-7-111-31493-6	金融市场与金融机构基础（第4版）（法博齐）（2010年）	外版	79
金融市场学	978-7-111-32325-9	金融市场与金融机构基础（英文版·第4版）（法博齐）（2010年）	外版	78
金融市场学	978-7-111-26841-3	现代金融市场：价格、收益及风险分析（布莱克威尔）（2009年）	外版	58
金融风险管理	978-7-111-30699-3	风险管理与金融机构（第2版）（约翰·赫尔）（2010年）	外版	56
金融风险管理	978-7-111-35863-3	风险管理与金融机构（英文版·第2版）（约翰·赫尔）（2011年）	外版	69
兼并、收购与公司重组	978-7-111-35538-0	兼并、收购和公司重组（第2版）（阿扎克）（2011年）	外版	62
国际金融案例	978-7-111-22673-4	国际金融案例（德赛）（2007年）	外版	48
（证券）投资学习题	978-7-111-29851-9	投资学习题集（第7版）（博迪）（2010年）	外版	35
（证券）投资学	978-7-111-26944-1	投资学（第7版）（博迪）（2009年）	外版	89
（证券）投资学	即将出版	投资学（第9版）（博迪）（2012年）	外版	89
（证券）投资学	978-7-111-27013-3	投资学（英文版·第8版）（博迪）（2009年）	外版	108
（证券）投资学	即将出版	投资学（英文版·第9版）（博迪）（2012年）	外版	108
（证券）投资学	978-7-111-23031-1	投资学：分析与管理（第10版）（琼斯）（2008年）	外版	69
（证券）投资学	978-7-111-30587-3	投资学：以Excel为分析工具（第3版）（霍顿）（2010年）	外版	38
（证券）投资学	978-7-111-30395-4	投资学基础（英文版·第9版）（梅奥）（2010年）	外版	82
行为金融学	978-7-111-31106-5	行为金融学（饶育蕾）（2010年）	本版	36
投资银行学	978-7-111-33063-9	投资银行：理论与案例（马晓军）（2011年）	本版	32
投资银行学	978-7-111-24899-6	投资银行学（张志元）（2009年）	本版	36
投资管理	978-7-111-29977-6	投资管理（李学峰）（2010年）	本版	36
投资管理	978-7-111-21926-0	投资管理（庄新田）（2007年）	本版	32
商业银行经营管理学	978-7-111-19856-5	商业银行经营管理学（刘毅）（2006年）	本版	32
金融学（货币银行学）	978-7-111-20910-2	货币银行学（精品课）（钱水土）（2007年）	本版	32
金融学（货币银行学）	978-7-111-38159-4	金融学（陈伟鸿）（2012年）	本版	35
金融学（货币银行学）	978-7-111-30153-0	金融学（精品课）（董金玲）（2010年）	本版	30
金融学（货币银行学）	978-7-111-35641-7	金融学概论（精品课）（丁志国）（2011年）	本版	48
金融学（货币银行学）	978-7-111-35022-4	金融学概论（精品课）（茆训诚）（2011年）	本版	42
金融投资专业英语	978-7-111-26455-2	金融投资专业英语（杜朝运）（2009年）	本版	25
金融市场学	978-7-111-28226-6	金融市场学（韩国文）（2009年）	本版	48
兼并、收购与公司重组	978-7-111-34929-7	并购与重组：跨世纪的十大经典重组（毛蕴诗）（2011年）	本版	38
国际结算	978-7-111-37202-8	国际结算（国家级双语示范课程）（邵新力）（2012年）	本版	39
国际结算	978-7-111-34376-9	国际结算（徐进亮）（2011年）	本版	38
风险投资学	978-7-111-29299-9	风险投资与创业板（房四海）（2010年）	本版	38
（证券）投资学	978-7-111-29863-2	投资学（朱相平）（2010年）	本版	38
（证券）投资学	978-7-111-35146-7	投资学原理及应用（贺显南）（2011年）	本版	29
（证券）投资学	978-7-111-21343-7	证券投资学（葛红玲）（2007年）	本版	35
（证券）投资学	978-7-111-23293-3	证券投资学原理（精品课）（韩德宗）（2008年）	本版	36

CFA特许金融分析师考试系列用书

机械工业出版社华章公司非常荣幸地为您提供CFA全球考试教材的中文版，本系列图书共包括11本，其中8本为考试用书，另外3本为推荐用书。全系列图书与CFA考试科目全面对接，涵盖CFA考试I II III 三级知识体系，原汁原味，帮助读者快速、整体掌握每门考试科目的核心知识要点。同时本系列丛书也是金融从业人员与管理者专业知识必备图书，帮你快速汲取国际先进经验，自如应对日益变化的金融市场。



CFA考试科目	书号	对应教材	备注	定
Quantitative Methods	978-7-111-38802-9已出版	定量投资分析（CFA协会投资系列）	考试用书	9
Financial Reporting and Analysis	2013年即将出版	国际财务报表分析（CFA协会投资系列）	考试用书	9
Corporate Finance	2013年即将出版	公司理财：实用方法（CFA协会投资系列）	考试用书	9
Equity Investments	978-7-111-38805-0已出版	股权资产估值（CFA协会投资系列）	考试用书	9
Fixed Income	2013年即将出版	固定收益证券分析（CFA协会投资系列）	考试用书	9
Portfolio Management and Wealth	978-7-111-38719-0已出版	投资组合管理（CFA协会投资系列）	考试用书	14
	2013年即将出版	财富管理（CFA协会投资系列）	考试用书	6
	2012即将出版	投资绩效测评（CFA协会投资系列）	考试用书	6
Alternative Investments	2012即将出版	投资学（博迪）	推荐用书	9
Derivatives	978-7-111-35821-3 已出版	期权、期货及其他衍生品（赫尔）	推荐用书	9
Economics	978-7-111-28088-0已出版	经济学：私人与公共选择（索贝尔）	推荐用书	7

订购热线：010-68995261 010-88379628

CFA全球考试教材中文版

- 原汁原味，全面系统，助您顺利通过CFA考试。
- 全面对接CFA考试科目，涵盖I、II、III 三级知识体系，迅速掌握每门考试科目的核心知识要点。
- 金融从业人员与管理者专业知识必备图书，快速汲取国际先进经验，自如应对日益变化的金融市场。

定量投资分析 (原书第2版)

Quantitative Investment Analysis (Second Edition)

对于当今金融市场中从事研究工作的每一位投资者来说，《定量投资分析》是一本不可不读的必备书籍。本书作者以简单易懂和具有说服力的语言为读者介绍了许多常用的定量投资方法。本书应该成为大家书桌上的必备读物和参考书籍。

格雷 P. 布里森 特许金融分析师，GP布里森投资公司主席

第2版的《定量投资分析》为投资从业者利用严格规范的定量工具和分析方法来获取利润的行为提供了坚实的理论基础。本书中大量关于定量投资分析的实例更是提升了自学者对本书学习的热情。

布莱恩·辛格 特许金融分析师，瑞银全球资产管理部

《定量投资分析》对于投资领域的核心内容予以了关注。本书语言清晰、结构合理、内容翔实，为读者理解投资组合管理和资产定价中统计学和经济学的知识进行了完美的论述。本书同样也对当今广泛使用的主要资产定价模型予以了详尽的介绍。

戴维·布莱克威尔 海德共和银行

詹姆斯 W. 阿斯頓 得克萨斯A&M大学迈斯商学院金融学教授

客服热线：
(010) 68995261, 88361066
购书热线：
(010) 68326294, 88379649, 68995259
投稿热线：
(010) 88379007
读者信箱：
hzjg@hzbook.com

华章网站 <http://www.hzbook.com>

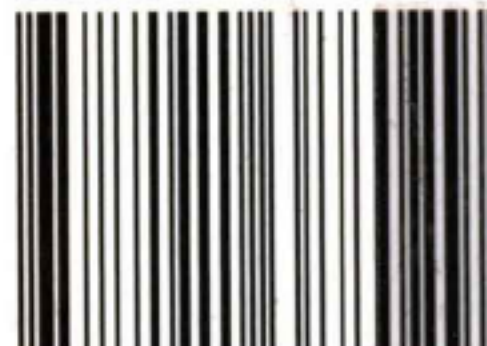
网上购书：www.china-pub.com

 **WILEY**
Publishers Since 1807
www.wiley.com



上架指导：金融投资

ISBN 978-7-111-38802-9



9 787111 388029

定价：99.00元